



Réseaux et séquents ordonnés

Christian Retoré

► To cite this version:

Christian Retoré. Réseaux et séquents ordonnés. Mathématiques [math]. Université Paris-Diderot - Paris VII, 1993. Français. NNT : . tel-00585634

HAL Id: tel-00585634

<https://theses.hal.science/tel-00585634>

Submitted on 13 Apr 2011

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.

UNIVERSITE PARIS 7
THESE DE DOCTORAT
Spécialité: MATHEMATIQUES

Présentée par: Christian RETORÉ

Sujet de la thèse:

Réseaux et Séquents Ordonnés.

Soutenue le 26 Février 1993

JURY: Jean-Yves GIRARD [Directeur de Recherche]
 Jean-Louis KRIVINE [Président du Jury]
 Jean-Jacques LOEB
 Jacques VAN DE WIELE
 Jacqueline VAUZEILLES
 Gilles ZEMOR

REMERCIEMENTS

C'est une grande chance pour moi d'avoir eu pour directeur de recherches un mathématicien aussi créatif que Jean-Yves Girard: il est en effet à l'origine de tous les aspects du sujet ici abordé. Je le remercie avant tout de ce sujet qu'il m'a confié, et de l'enthousiasme pour la logique linéaire qu'il sait si bien faire partager. Sa perspicacité mathématique et ses goûts très sûrs m'ont préservé des "atrocités logiques" qui guettent le novice.

Je suis très reconnaissant à Jean-Louis Krivine de m'avoir fait l'honneur de présider mon jury. Ses travaux, cours et séminaires sont un exemple qui donne envie de devenir mathématicien.

Je remercie beaucoup Jean-Jacques Loeb qui a bien voulu participer au jury. Sa curiosité, l'étendue de son savoir et sa gentillesse en font un mathématicien irremplaçable. Si plus de collègues partageaient son ouverture d'esprit, il n'y aurait assurément pas de ghettos mathématiques, tels la logique, la combinatoire ou les catégories.

C'est un honneur pour moi qu'Anne Troelstra ait apprécié ce travail, attendu que l'intuitionnisme est une conception des mathématiques qui m'a toujours attiré.

Je remercie beaucoup Jacqueline Vauzeilles que j'ai rencontrée alors que je rédigeais ce travail. L'intérêt qu'elle y a porté, sa compréhension du sujet, et la gentillesse de ses encouragements m'ont beaucoup aidé à en venir à bout.

Un immense merci à Jacques Van de Wiele: sa disponibilité sans pareille et sa générosité m'ont souvent permis de profiter de son excellente connaissance de la logique et de la profondeur de sa réflexion. Ce travail doit énormément à la pertinence de ses suggestions, et à ses encouragements constants.

Je remercie amicalement Gilles Zémor, dont le vaste savoir et la vivacité mathématique m'ont souvent éclairé: la partie concernant les graphes doit beaucoup à nos discussions où il a vaillamment enduré de très pénibles versions préliminaires.

Je remercie de même Arnaud Fleury, avec qui j'ai commencé ce travail. Ses facilités mathématiques nous ont permis d'accéder rapidement au cœur du sujet. Je lui souhaite de mener rapidement à bien l'ambitieuse thèse qu'il prépare.

Mille et un mercis à Catherine Amparo Gourion, qui a tant contribué à cette thèse: la rédaction du premier chapitre, les dessins, le temps dont j'ai disposé alors qu'elle gardait nos enfants... Par chance, sa thèse prochaine devrait me permettre de m'acquitter de cette dette.

Je remercie mes laboratoires d'accueil où j'ai pu travailler dans de bonnes conditions:

- * l'équipe de logique de l'université Paris 7 (où l'irremplaçable Madame Orieux conjugue si bien amabilité et efficacité)
- * the "Computing Theory and Formal Methods" team of the Imperial College (London) where Samson Abramsky nicely invited me: I thus benefited from the best introduction to the computational aspects of linear logic.
- * le dynamique et chaleureux département de mathématiques de l'université d'Angers et plus particulièrement les logiciens angevins — M. Brestovski, J.-L. Duret, D. Glushankof, F. Lucas

Je remercie enfin de leur soutien amis et collègues: P. Ageron, M.-F. Allain, J. Alexander, F. Ducrot, E. Duquesne, M. El Amrani, J.-M. Granger, C. Lair, P. Malacaria, F. Métayer, H. Schellinx, Paul Taylor (who wrote the LaTeX package I use for proof trees) et Michel Parigot (que j'ai trop souvent dérangé)... ..ainsi que tous ceux que j'ai assurément oubliés!

Je dédie cette thèse à mes grands-parents, Georges et Isabelle Retoré, Gabriel et Mireille Allais, et à ce monde paysan qui, je le crains, disparaîtra avec eux.

Les cerveaux des oiseaux, des vaches, des personnes
N'émettent c'est bizarre aucun son dans les champs

Jacques RÉDA

Introduction

Le cadre général de ce travail en théorie de la démonstration est l'étude des preuves considérées comme des programmes [Par88, GLT88, Kri90, How80], et plus particulièrement celle des preuves de la logique linéaire [Gir90, Gir87a, Tro92]. Les techniques utilisées étant souvent de nature combinatoire, on se référera éventuellement aux classiques du genre tels [Ber73, Ber79, GW75, Bol79].

L'essentiel de cette thèse est consacrée à l'étude combinatoire des réseaux de preuve de la logique linéaire multiplicative, tout d'abord pour le calcul avec la règle de MÉLANGE, puis (et c'est là le centre de notre travail) pour un calcul nouveau portant sur des multi-ensembles ordonnés de formules, qui étend le premier et peut se concevoir comme un modèle du parallélisme. Dans ce calcul cohabitent des connecteurs généralisés définis par des ordres, les connecteurs commutatifs et associatifs usuels **par** ($\wp$) et **tenseur** ($\otimes$), et un autre baptisé **précède** ($<$), qui possède les étonnantes propriétés suivantes: ce connecteur est associatif, non-commutatif, autodual ($(A < B)^\perp = (A^\perp < B^\perp)$) et il est situé, pour l'ordre défini sur les formules par l'implication linéaire, entre le **tenseur** et le **par**.

The general framework of this proof-theoretical work is the study of proofs-as-programs [Par88, GLT88, Kri90, How80], and more precisely the proof-theoretical study of linear logic [Gir90, Gir87a, Tro92]. We often make use of combinatorial techniques for which the reader can refer to the classical books in the area like [Ber73, Ber79, GW75, Bol79].

*This thesis is mainly concerned with proofnets of multiplicative linear logic, firstly for a calculus with the MIX-rule, secondly (and this the heart of our work) for a new calculus dealing with partially ordered multisets of formulae, which extends the first one and which may be viewed as a model for concurrency. The latter involves generalised connectives defined by orders, the usual associative and commutative **par** ($\wp$) and **tenseur** ($\otimes$), and a third one christened **précède** ($<$), which enjoys the following unusual properties: this connective is associative, non-commutative, self-dual ($(A < B)^\perp = (A^\perp < B^\perp)$), and is in between **par** and **tenseur** for the partial order defined on formulae by linear implication.*

Cette restriction à un calcul purement multiplicatif est juste un moyen de mieux cerner les problèmes rencontrés; nous évoquons dans la conclusion comment incorporer à notre calcul les modalités et les quantificateurs (ce qui est trivial), pour ainsi obtenir un système d'un pouvoir d'expression raisonnable (contenant la logique intuitionniste du second ordre).

Le calcul ordonné est au confluent de deux courants de la logique linéaire: l'étude des connecteurs généralisés initiée dans les articles [Gir87b] et [DR90], quoique nos connecteurs généralisés soient d'une autre nature, et l'étude de la logique linéaire non-commutative [Gir87a, Yet90], qui se propose d'affaiblir la règle d'échange — quoique là encore notre point de vue soit radicalement différent. Il permet de plus de traiter les éléments neutres et l'affaiblissement tout en évitant certains inconvénients de la règle de mélange.

Pour pouvoir appeler calcul ce calcul ordonné, nous avons commencé par établir les propriétés suivantes, qui constituent à notre avis le minimum pour une telle appellation:

★ une sémantique des preuves non-dégénérée (ou modèle de Heyting) qui garantisse le bien-fondé et le caractère constructif de ce calcul, et qui est ici fournie par les espaces cohérents,

★ un calcul en réseaux, c'est-à-dire un critère simple qui permette de reconnaître les preuves des structures ressemblantes, et qui soit préservé par élimination des coupures,

★ un calcul des séquents tel que toute preuve séquentielle se traduise en un réseau correct

This restriction to the plain multiplicative calculus is just a way to outline the encountered difficulties; in the conclusion we sketch how to deal with modalities and quantifiers (this is straight forward) in order to obtain a calculus of a sensible expressive power (including the second order intuitionistic logic).

Two trends of linear logic merge in the ordered calculus: the study of generalised connectives started by [Gir87b] and [DR90], even though the nature of our connectives is completely different, and the study of non-commutative linear logic [Gir87a, Yet90] which limits the exchange rule — even though our approach once again differs from theirs. Moreover this calculus provides a proofnet calculus for units and weakening which avoids some drawbacks of the mix rule.

In order to call calculus the ordered calculus, we firstly established the following properties, which, according to us, are a minimum for deserving such an appellation:

★ *a non-degenerated denotational semantics or Heyting model, here provided by coherence spaces, which ensures the soundness and the constructiveness of the calculus,*

★ *a proofnet syntax, i.e. a simple criterion, preserved by cut-elimination, which distinguishes between proofs and closely related objects,*

★ *a sequent calculus such that any sequential proof is translated into a correct proofnet.*

Résumons brièvement notre apport aux deux courants sus-mentionnés.

Concernant la logique linéaire non-commutative, l'étude s'est essentiellement restreinte au calcul des séquents (sans modalités ni quantificateurs), à celle des valeurs de vérité ou phases (avec les modalités mais pas de quantificateurs) tandis que:

- * nous disposons d'un calcul en réseaux (et en séquents) qui s'étend sans peine aux modalités et quantificateurs,

- * nous disposons d'une sémantique des preuves pour ce calcul, où le connecteur non-commutatif s'interprète par un objet non-isomorphe à son symétrique

- * dans ce système, les connecteurs usuels et commutatifs cohabitent avec le connecteur non-commutatif

Concernant les connecteurs généralisés, nous avons réussi à leur donner:

- * une sémantique des preuves (fournie par les espaces cohérents); cette question avait été posée pour les connecteurs généralisés de [Gir87b, DR90], mais le simple calcul des espaces cohérents correspondant n'admet, aujourd'hui encore qu'une réponse partielle (mais très élégante) en termes de jeux due à [vdW89];

- * un calcul des séquents traitant de ces connecteurs, dont nous avons réussi à prouver qu'il était le plus riche à ne donner que des réseaux corrects (par une étude de l'orthogonalité des préordres, i.e. du type des modules);

Let us briefly summarize our contribution to the two quoted trends.

Concerning non-commutative linear logic the study was essentially limited to the sequent calculi (without modalities and quantifiers), to the truth values or phases (with modalities but no quantifiers) while:

- * we have a proofnet calculus (and a sequential calculus) which may be easily extended to modalities and quantifiers*

- * we have a proof semantics for this calculus, where the non-commutative connective is interpreted as an object non-isomorphic to its symmetric*

- * in this system, the usual commutative connectives and the non-commutative one live under the same roof*

For generalised connectives we manage to give them:

- * a proof semantics (provided by coherence spaces); this question had already been pointed out for the generalised connectives of [Gir87b, DR90], and [vdW89] has given an elegant but partial solution in terms of games;*

- * a sequent calculus dealing with such connectives, which is proved to be the largest to give correct proofnets (by a study of orthogonality of preorders i.e. of the types of the proof modules);*

* une quasi-preuve de la séquentialisation, que nous mentionnons, bien qu'un maillon de la preuve nous fasse encore défaut; en effet ce serait tout à fait nouveau pour des connecteurs généralisés. On montre aussi que le calcul des séquents introduit est le seul possible pour avoir ce théorème de séquentialisation.

* une interprétation intuitive qui s'insère dans le cadre très général de l'isomorphisme de Curry-Howard entre preuves et programmes, et de la logique linéaire comme une logique des actions.

Ce dernier point est bien sur la nouveauté essentielle, et vient du fait que l'ordre porte aussi sur les coupures, c'est à dire sur les calculs à effectuer. Une stratégie d'évaluation se trouve ainsi décrite, et l'ordre étant partiel, les calculs minimaux sont à effectuer en parallèle, tandis que certains calculs sont à effectuer avant d'autres: cela donne lieu à des optimisations de stratégies, et à une interprétation en termes de parallélisme que nous décrivons ci-après.

** an almost complete proof of sequentialisation; although a lemma remains unproved, we include its formulation and expose the established lemmas since it would be a completely new result for generalised connectives. We show that the sequent calculus we give is the only possibility to have the sequentialisation theorem.*

** an intuitive meaning for these connectives which fits into the Curry-Howard isomorphism between proofs and programs, as well as into the interpretation of linear logic as a logic of actions.*

This latter point is of course the major novelty, and comes from the fact the order also concerns the cuts i.e. the computations to be performed. A computational strategy is thus described, and as the order is a partial order, the minimal computations are to be computed in parallel, while some computations are to be performed before others: this give rise to an interpretation as a model for concurrency to be below described.

§A. Résumé

Commençons par la partie B, qui est la plus standard et aussi l'origine de ce travail:

Le calcul avec la règle de mélange (Partie B)

Comme ce calcul ordonné est une extension du calcul avec la règle de mélange, nous traitons préalablement de ce dernier, qui a son intérêt propre: il permet, en effet, de se passer des boîtes $\perp$ de [Gir87a] (qui sont évidemment les mêmes que les boîtes affaiblissement). Un examen minutieux des cas de base de l'élimination des coupures de [Gir87a] montre que sans ces boîtes, le calcul avec les modalités et les quantificateurs conflue.

Néanmoins, si l'on ne restreint pas la règle de mélange, on a: $1 = \perp$, ce qui signifie qu'une formule $!A$ se comporte comme une formule affaiblie. On obtient alors des preuves de conclusions $!A, !B$, qui ne se trouvent pas dans le calcul usuel.

Nous remédions à ce problème de la manière suivante: on indexe les séquents par des entiers, et l'on demande que l'index final soit 0. Il est inutile d'ajouter cette information aux réseaux, car il se déduit du nombre de composantes connexes de tout graphe de correction. Cet index ajouté est un invariant de l'élimination des coupures, et fournit une sémantique des phases.

On obtient alors un calcul où $1 \neq \perp$ et où la prouvabilité des formules avec $1, \perp, \otimes, \wp$ est décidable en temps polynomial, ce qui est assurément préférable à la NP-complétude de ce problème pour le calcul habituel [LW92].

§A. Summary

Let us start by part B, which is the more standard one and the starting point of this work:

The mix-calculus (Part B)

Since the ordered calculus extends the mix-calculus we firstly deal with the mix calculus, which also has its own interest: it is a way to get rid of the $\perp$ -boxes of [Gir87a] (which are clearly the same as the weakening boxes). A close examination of the basic cases of cut-elimination of [Gir87a] shows that without these boxes the calculus with modalities and quantifiers is confluent. Nevertheless if we do not limit this mix-rule, we obtain $1 = \perp$, and this means that a $!A$ formula behaves like weakened formula. We there fore obtain proofs with conclusions $!A, !B$ contrary to what happens in the usual calculus.

We avoid this drawback as follows: we index sequents by integers, and we ask for the final index to be 0. There is no need to add anything to the proofnet calculus as it may be computed from the invariant number of connected components of any correction graph of the proofnet. This index is preserved by cut-elimination and gives a phase semantics for this calculus.

We thus obtain a calculus where $1 \neq \perp$ and where the provability of formulae written with $1, \perp, \otimes, \wp$ is decidable in polynomial time, which clearly improve the usual case for which it is NP-complete [LW92].

Notons que notre critère pour ce calcul: "sans cycle et (nombre d'affaiblissement + 1) = (nombre de composante connexes)" se généralise récursivement au calcul plein de la manière suivante: un réseau est correct si et seulement si l'intérieur de toute boîte maximale est un réseau correct et si le réseau obtenu en remplaçant toute boîte maximale par un axiome n -aire est correct. On obtient ainsi un calcul correspondant à [Abr91] dans lequel on sait reconnaître les réseaux ("proofs") des pré-réseaux ("proof expressions").

Mais notre principal intérêt pour ce calcul est que ces réseaux sont le calcul sous-jacent du calcul ordonné. Ayant du affiner notre connaissance de ces réseaux pour l'étude du calcul ordonné, nous avons été conduit à introduire la notion d'agrégat, dont on établit clairement qu'ils ont les mêmes propriétés combinatoires que les réseaux tout en étant plus maniables et plus élégants.

Combinatoire (Partie A)

Agrégats (Chapitre 2)

Ce sont des graphes dont les arêtes sont colorées tels que pour toute couleur le sous-graphe de cette couleur soit biparti-complet. On peut aussi les voir comme une généralisation particulière des graphes, différente des hypergraphes: une arête est une paire d'ensembles finis de sommets (au lieu d'une paire de sommets). L'analogie se poursuit ainsi: on s'intéresse aux chemins dits bigarrés, ceux qui n'empruntent pas deux arêtes de la même couleur — qui correspondent aux chemins réalisables des réseaux (selon la terminologie de [Dan90]); ce sont l'analogue des chemins simples (qui n'utilisent pas deux fois une même arête). On s'intéresse ensuite aux couleurs scindantes i.e. aux couleurs dont la suppression des arêtes déconnecte totalement les deux parties du graphe biparti-complet correspondant; ce sont donc l'analogue des isthmes.

Notice our criterion "without cycle and (number of weakening +1) = (number of connected components)" may easily be recursively generalised to the full calculus this way: a proofnet if and only if the inside of any maximal box is correct, and if the proofnet obtained by replacing maximal boxes by n -ary axioms is correct. We thus obtain a calculus corresponding to [Abr91] in which we are able to recognize proofnets ("proofs") from preproofnets ("proof expressions").

But our main interest in this calculus is that it is the underlying calculus of the ordered calculus. As we had to improve our knowledge on these proofnets to the study the ordered ones we were lead to introduce the aggregates notion which we show to be the same combinatorial objects as proofnets, although they are more manageable and elegant.

Combinatorics (Part A)

Agregates (Chapter 2)

They are edge coloured graphs such that for any colour the corresponding subgraph is complete bipartite. They may also be viewed as a particular generalisation of graphs, different from hypergraphs: an edge is a pair of disjoint finite sets of vertices (instead of a pair of distinct vertices). The analogy goes as follows: we are interested in variegated paths, i.e. path not using twice the same colour corresponding to the feasible paths of a proofnet (according to the terminology of [Dan90]); they are the analogue of simple paths. Then we look at splitting colours, i.e. colours whose removal totally disconnects the two parts of the corresponding complete bipartite graph; they are the analogous of bridges.

On démontre alors le théorème suivant: *tout point d'un agrégat est relié par un chemin bigarré à un point d'une couleur scindante, ou à un point d'un cycle bigarré.* Les réseaux se traduisant fidèlement (pour les chemins bigarrés) de deux manières en des agrégats sans cycle bigarré, ce théorème a pour corollaires triviaux l'existence d'un *par* scindant dans un réseau, et l'existence d'un *tenseur* scindant dans un réseau dont les conclusions sont toutes des *tenseur*; il permet de plus d'établir l'existence de certains chemins dans un réseau, qui sont utiles à l'étude des réseaux ordonnés. On montre enfin que le nombre de composantes connexes de tous les sous-graphes bigarrés maximaux est le même — cela nous permettra de traiter les neutres au chapitre 5.

Nous avons pensé que cette notion d'agrégat — à la différence des réseaux — et le résultat les concernant, pouvaient avoir un intérêt combinatoire intrinsèque: c'est ce qui explique qu'on en ait fait un chapitre indépendant.

Ordres et Orthogonalité (Chapitre 1)

Un autre aspect purement combinatoire de notre travail et nécessaire à l'étude du calcul ordonné est l'étude des ordres finis et de l'orthogonalité entre les relations binaires non symétriques.

On définit, dans ce chapitre, deux opérations de contraction sur les ordres qui correspondent respectivement à la formation de $A \wp B$ et de $A < B$ lorsque les formules A et B font partie d'un ensemble ordonné de formules et de coupures. On définit aussi les deux opérations inverses qui correspondent au remplacement d'une coupure par deux plus petites dans un ensemble ordonné de coupures et de formules, la première correspondant à une coupure $\wp/\otimes$, la seconde à une coupure $</<$.

We then prove the following theorem: any vertex of an aggregate is connected via a variegated path to a vertex of a splitting colour, or to a vertex of a variegated cycle. Proofnets are faithfully (according to variegated paths) translated in two ways into aggregates, this theorem has as straight forward corollaries the existence of splitting *par*, and the existence of a splitting *tenseur* in a proofnet whose conclusions are all $\otimes$ -formulae.; this theorem also establishes the existence of particular paths in a proofnet, which are usefull in the study of ordered proofnets. We then proof that the number of connected components of any maximal variegated subgraph is the same, and this enables us to deal with units in chapter 5.

We think that this aggregate notion — as opposed to proofnet — have an intrinsic combinatorial interest; that is the reason why we made an independant chapter of their study.

Orders and Orthogonality (Chapter 1)

Another purely combinatorial aspect of our work, necessary for the study of the ordered calculus is the study of finite orders and of orthogonality of non symmetric binary relations.

We define in this chapter, two operations on orders of contraction which corresponds respectively to the formation of the formulae $A \wp B$ and $A < B$ when the formulae A and B belong to an ordered (multi)set of formulae and cuts. We also define the two converse operations which correspond to the replacement of a cut by two smaller cuts within an ordered set of formulae and cuts, in the first case for a $\wp/\otimes$ cut and in the second for a $</<$ cut.

Les contractions identifient soit deux points indiscernables dans l'ordre, soit deux points dont le premier est l'unique prédécesseur du second, et le second l'unique successeur du premier. On démontre alors le théorème suivant: un ordre fini se contracte en un point si et seulement s'il satisfait la condition globale suivante: $(A < A' \text{ et } A < B \text{ et } B' < B) \Rightarrow (A' \leq B \text{ ou } A \leq B' \text{ ou } B' \leq A')$.

Ce théorème permet de caractériser les connecteurs généralisés définis par un ordre qui sont équivalents à une formule; les contractions de même espèce étant confluentes, cette formule est unique modulo l'associativité et la commutativité du *par* et l'associativité du *précède*.

On définit ensuite l'orthogonalité entre relations: deux relations sont orthogonales lorsque leur composée est sans "vrai" circuit (i.e. $(x, y) \in c \Rightarrow (y, x) \notin c$). On caractérise alors l'ordre le plus riche qui ne crée pas de circuit (cycle orienté) non-plat lorsqu'on forme un TENSEUR ou un MÉLANGE dans les réseaux ordonnés. Cela permet aussi de définir les modules du calcul ordonnés.

Réseaux et séquents ordonnés (Partie C)

On commence par présenter l'espace cohérent associé au *précède* et le produit ordonné d'espaces cohérents, puis on définit un calcul en réseaux qui correspond à cette sémantique. Enfin on présente le calcul des séquents correspondant, et on étudie la séquentialisation.

Des espaces cohérents au connecteur *précède* (Chapitre 6)

En appelant connecteur multiplicatif les lois de composition sur "*∪*", "*=*", "*∧*" qui sont croissantes (pour l'ordre "*∪*" < "*=*" < "*∧*"), on voit qu'il n'y a que trois connecteurs multiplicatifs: le *par*, le *tenseur* et le *précède*.

Contractions identify either two indiscernible points, or two points the first of them being the only predecessor of the second one, and the second one the only successor of the first one. We then prove the following theorem: a finite order contracts to a vertex if and only if it enjoys the following global property: $(A < A' \text{ and } A < B \text{ and } B' < B) \Rightarrow (A' \leq B \text{ or } A \leq B' \text{ or } B' \leq A')$.

This theorem provides a characterisation of generalised connectives which are equivalent to a single formula; because of the confluence of contractions of the same kind, the obtained formula is unique modulo the associativity and commutativity of *par* and the associativity of *précède*.

We then define orthogonality between relations: two relations are orthogonal when their composition has no "real" circuit (i.e. $(x, y) \in c \Rightarrow (y, x) \notin c$). We thus characterize the largest order which does not creates any circuit (directed cycle) when we use a TENSEUR or MÉLANGE rule in the ordered proofnets. This also leads to the notion of modules for ordered proofnets.

Ordered sequents and proofnets (Part C)

We firstly introduce the coherence space corresponding to *précède* and the ordered product of coherence spaces, then we define a proofnet syntax corresponding to this semantics. Finally we give the sequent calculus and study the sequentialisation.

From coherence spaces to the *précède* connective (Chapter 6)

If we call multiplicative connective any binary operation on "*∪*", "*=*", "*∧*" which are increasing (with respect to the order "*∪*" < "*=*" < "*∧*"), only three multiplicative connectives exist: the *par*, the *tenseur* and the *précède*.

On vérifie alors que ce connecteur *précède* est non-commutatif, associatif, et auto-dual, et situé entre les connecteurs *par* et *tenseur* pour l'implication linéaire. En généralisant la définition de la cohérence selon le *précède*, on définit le produit d'une famille d'espaces cohérents A_i ordonnée par i : $(x_1, \dots, x_n) \wedge (x'_1, \dots, x'_n)$ ssi $\exists i [x_i \wedge x'_i \text{ et } \forall j > i[i] x_j = x'_j]$.

Les réseaux ordonnés (Chapitre 7)

On définit un calcul correspondant à la sémantique définie plus haut. Ces réseaux généralisent très naturellement les réseaux de la seconde partie. Le lien correspondant au connecteur *précède* est le même que celui du lien *par*, auquel on ajoute un arc de la prémisses A vers la prémisses B si la conclusion est $A < B$. L'ordre sur l'ensemble des conclusions et des coupures est représenté par un ensemble d'arcs: on a un arc de F vers G si et seulement si $F < G$ — où F et G sont des coupures ou des conclusions. Le critère est le suivant: *tout graphe de correction ne contient pas de circuit (cycle orienté)*.

Les étapes élémentaires d'élimination sont semblables à celles du calcul habituel. Néanmoins on doit préciser quel est l'ordre sur les conclusions et coupures après une telle étape: dans le cas d'une coupure entre un *par* et un *tenseur* les deux coupures créées sont équivalentes dans le nouvel ordre (et elles occupent dans l'ordre la place de celle qu'on vient de supprimer) tandis que dans le cas d'une coupure entre deux *précède* l'une est des deux coupures créées *précède* l'autre (et elles occupent la place de celle que l'on vient de supprimer). On montre alors que le critère de correction est préservé par élimination des coupures, et que toute preuve de $A_1, \dots, A_n$ comportant les coupures $c_1, \dots, c_n$ dans un ordre i portant, et sur les A_i , et sur les c_i se normalise en une preuve de $A_1, \dots, A_n$ dans l'ordre i restreint à $A_1, \dots, A_n$. Cette normalisation est évidemment forte et confluyente.

We then check that this latest *précède* connective is non-commutative, associative, and self-dual; moreover it is in between *tenseur* and *par* for the order corresponding to linear implication. We then generalise the coherence with respect to *précède* to define the product of a family A_i of coherence spaces ordered by i : $(x_1, \dots, x_n) \wedge (x'_1, \dots, x'_n)$ iff $\exists i [x_i \wedge x'_i \text{ et } \forall j > i[i] x_j = x'_j]$.

Ordered proofnets (Chapter 7)

We here define a calculus corresponding to the semantics of the previous chapter. These proofnets are a straight forward generalisation of the proofnets of part B. The link corresponding to the *précède* connective is the *par* link leaded with an arc from A to B if the conclusion is $A < B$. The order on conclusions and cuts is encoded by arcs: there is an arc from F to G iff $F < G$ — F and G being conclusions or cuts. The correctness criterion is the following: any correction graph contains no circuit (directed cycle)

The cut-elimination steps look like the usual ones, but one must specify the order on the reduced proofnet: if the cut-elimination step concerns a *par* and a *tenseur* formula, the two smaller cuts are equivalent in the new order (and take the place of the removed cut); if the cut-elimination step concerns two *précède* formulae, the two smaller cuts are lower-equivalent in the new order (and take the place of the removed cut). We then prove that the criterion is preserved by cut-elimination and that any proof having $A_1, \dots, A_n$ as conclusions, $c_1, \dots, c_n$ as cuts, i as order concerning both the A_i , and the c_i reduces to a cut-free proof of $A_1, \dots, A_n$ in the order i restricted to $A_1, \dots, A_n$. This normalisation is obviously strong and confluent.

On définit ici la sémantique cohérente d'une preuve par une généralisation de la méthode des expériences [Gir87a]. On montre que deux expériences d'un réseau correct sont toujours compatibles (modulo le produit ordonné des espaces cohérents). On montre alors que les réseaux normaux ont une sémantique non-triviale, et que l'élimination des coupures préserve la sémantique cohérente. Ainsi tout réseau ordonné a-t-il une sémantique cohérente non-triviale. Ceci montre que notre calcul ordonné représente suffisamment de fonctions, ou de programmes.

On démontre ensuite que dans un réseau ordonné il est toujours possible de transformer les liens *précède* en liens *tenseur* ou par de manière à obtenir un réseau habituel (c'est trivial pour les réseaux sans coupure et plus délicat pour ceux avec coupures: il faut utiliser des résultats des chapitres 2 et 4). Cela montre que ce calcul est une extension conservative du calcul de la seconde partie non seulement par rapport aux preuves, mais aussi par rapport à la dynamique des calculs.

On démontre aussi que les formules $A_1, \dots, A_n$ dans l'ordre $\dot{}$ sont équivalentes à une formule $F(A_1, \dots, A_n)$ si et seulement si l'ordre $\dot{}$ est contractile, ce qu'on a défini et caractérisé au chapitre 1. Dans ce cas, la formule F ne contient que des *précède* et des *par* et est unique modulo l'associativité près de ces deux connecteurs, et la commutativité du second. On peut, d'après le chapitre 1, caractériser ces ordres.

Calcul des séquents ordonnés (Chapitre 8)

On définit ici les règles du calcul des séquents. On prend pour chaque règle la règle la plus libérale à ne créer que des réseaux corrects, ce que les chapitres 1 et 7 permettent de définir. On vérifie alors que ces règles sont interprétées par la sémantique cohérente. Le calcul choisi est un ainsi le plus libéral à ne donner que des réseaux corrects et à être interprété par la sémantique cohérente.

We here define the coherence semantics of a proof by a generalisation of the experiment method [Gir87a]. We prove that any two experiment of a correct proofnet are coherent, that normal proofnets have a non-trivial coherence semantics, and that the coherence semantics of a proofnet is preserved by cut-elimination. Thus any proofnet has a non-trivial coherence semantics. This shows that our ordered calculus contains enough functions or programs.

*We then prove that any *précède* connective of an ordered proof net may be transformed into a *tenseur* or *par* in order to obtain a usual proofnet (this is straight forward for cut-free proofnets and trickier for proofnet with cuts: one applies results from chapters 2 and 4). This shows that this calculus is a conservative extension of the calculus of the second part, not only with respect to proofs, but also with respect to the dynamics of computation.*

*We also prove that the formulae $A_1, \dots, A_n$ ordered by $\dot{}$ are equivalent to a single formula $F(A_1, \dots, A_n)$ if and only if the order $\dot{}$ is contractile, which we already defined and characterised in chapter 1. In this case, the formula F only contains *précède* and *par*, and is unique modulo the associativity of these two connectives, and the commutativity of the latest.*

The calculus of ordered sequents (Chapter 8)

We here define the rules of the sequent calculus. We take the most permissive rules which only lead to correct proofnets, and chapters 1 and 7 enables us to define such rules. We then check that these rules admit a coherence semantics interpretation. Therefore the obtained calculus is the largest to be sound with respect to both coherence semantics and proofnet syntaz.

Séquentialisation des réseaux ordonnés (Chapitre 9)

On donne ici une quasi-preuve de la séquentialisation, i.e. une preuve reposant sur un lemme de théorie des graphes dont la démonstration n'est pas encore achevée. On procède par la méthode du tenseur scindant. Il faut, après avoir supprimé les par et précède finaux augmenter l'ordre de manière à trouver une frontière correspondant à une règle TENSEUR ou MÉLANGE finale; des exemples, issus du chapitre 2, montrent alors qu'on ne peut se restreindre à ne rechercher qu'un seul type de frontière, c'est-à-dire que et que les règles MÉLANGE et TENSEUR ne commutent pas (sur les formules, elles commutent évidemment, mais pas sur les ordres).

Les frontières correspondant à des règles MÉLANGE et TENSEUR finales sont étudiées, et si aucune ne correspond à une règle alors nombre de chemins sont simultanément réalisables. Le lemme admis montre qu'alors il y aurait un circuit, ce qui est impossible dans un réseau correct.

D'autres exemples du même style montrent que si l'on restreint les règles du calcul des séquents, par exemple si l'ordre créé par ces deux règles TENSEUR et MÉLANGE est toujours d'un séquent prémisses vers l'autre, alors il existe des réseaux non-séquentialisables.

La preuve n'étant pas finie on envisage aussi ce qui se passerait si le résultat était faux: on aurait alors un calcul des réseaux qui satisferait toutes les propriétés habituelles mais qui n'aurait aucun calcul des séquents correspondant, puisque nous avons montré que le calcul des séquents proposé est le plus riche possible à ne donner que des réseaux corrects; ce serait là une situation complètement nouvelle.

Sequentialisation of ordered proofnets (Chapter 9)

We give here an almost complete proof of sequentialisation, i.e. a proof which relies on a graph theoretical lemma whose proof remains unfinished. We proceed by the splitting tensor method. After suppressing the final précède and par links, one has to find a border corresponding to a final MÉLANGE or TENSEUR rule. Examples coming from chapter 2, show that it is not enough to look for just one kind of border, i.e. the MÉLANGE et TENSEUR rules do not commute — with respect to formulae they obviously do, but not with respect to orders.

The borders corresponding to final MÉLANGE or TENSEUR rules are studied and we show that if no border corresponds to a final MÉLANGE or TENSEUR rule then numerous paths would be simultaneously feasible. The admitted lemma says that the existence of all these simultaneously feasible paths would imply a feasible circuit, and this is impossible in a correct proofnet.

Other examples of the same kind also shows that when the rules are restricted to only create order from one premisses sequent to the other, there exists non-sequentialisable proofnets.

As the proof is not finished we also consider what would happen if the sequentialisation should fail. In this case we would obtain a proofnet calculus with all the usual good properties, without any corresponding sequent calculus since we have proved that the given sequent calculus is the largest which give only correct proofnets; this would be a radically new situation.

§B. Interprétation Naïve du Calcul Ordonné comme Modèle du Parallélisme

On aura remarqué qu'il n'y a guère plus de rapport entre un de nos connecteurs et le "ou" usuel qu'il n'y en a entre un entier non-standard et le nombre 2: notre système ne mérite l'adjectif logique qu'en cela qu'il possède (plus ou moins) les mêmes propriétés que la logique intuitionniste ou classique.

Fort heureusement, l'informatique fournit un support relativement concret où interpréter ce calcul. Dans notre calcul ordonné, l'ordre porte, et sur les formules conclusions, et sur les coupures c'est-à-dire sur les calculs à effectuer. Chacun sait que l'élimination des coupures dans les réseaux de preuves est usuellement un calcul où tout ce qu'il est possible de faire en parallèle est effectivement représenté par des coupures pouvant être calculées en parallèle. Néanmoins ce parallélisme outrancier n'est pas toujours idéal: par exemple, une coupure sur un affaiblissement a tout intérêt, pour que le calcul termine au plus vite, à être effectuée en premier, puisque sa réduction fait disparaître tout une preuve qui contenait éventuellement des coupures.

Il est alors utile (et naturel) d'interpréter notre ordre sur les coupures comme un ordre dans lequel les réduire. On peut alors décider que lors d'un affaiblissement, la formule affaiblie soit minimale dans l'ordre: ainsi une coupure sur une telle formule sera t elle éliminée en premier¹. Ainsi cet ordre sur les coupures et conclusions s'interprète comme une stratégie qui est décrite dans la syntaxe même.

§B. A Naïve Interpretation of the Ordered Calculus as a Model for Concurrency

One easily checks that there are no more relationships between our connectives and the usual "or" than there are between a non-standard integer and the number 2: our system only deserves the adjective "logical" because it (more or less) enjoys the same properties as intuitionistic or classical logic.

Fortunately, computer science provides a relatively concrete medium to interpret this calculus. In this ordered calculus the order concerns both the formulae and the cuts i.e. the computations to be performed. It is well-known that cut-elimination on proofnets is a computation in which any possibly parallel computations are actually encoded by parallel computations. Nevertheless this excessive parallelism is not always absolute: for instance, a cut on a weakening should be performed first, if we want the computation to be as quick as possible, since its computation erases a part of the proof which possibly contains cuts.

Thus it is quite usefull (and natural) to interpret the order on the cuts as an order on the computations. We can now decide that a weakened formula be minimal in the order: this way a cut on such a formula will always be performed first². This order is therefore interpreted as a strategy which is described within the syntax itself.

¹en fait il vaut mieux placer cette formule affaiblie en dessous de toutes les formules de hauteur 1 de l'ordre, pour préserver l'associativité de l'affaiblissement

²Actually, we should rather put such a formulae below all formulae of height 1 in the order, to preserve the associativity of the weakening rule

Notons qu'il s'agit là de stratégie plutôt que de contraintes et de parallélisme, car le calcul aboutit toujours et notre calcul interprète naturellement une règle ENTROPIE qui oublie (en partie) l'ordre.

Par contre, en présence d'axiomes autres que ceux de la logique pure, alors cette règle qui permet d'oublier l'ordre doit être exclue. En effet, si ces axiomes extra-logiques expriment des contraintes destinées à un robot, tels "ouvrir la porte avant d'entrer", comme au cours de la preuve cet ordre se mêle aux autres ordres, il serait manifestement maladroit de permettre à un quelconque endroit de la preuve un affaiblissement de cet ordre.

Ainsi en présence d'axiomes extra-logique, cet ordre modélise-t-il la dépendance et l'indépendance temporelle des actions et leur indépendance, c'est-à-dire un calcul parallèle³. Son seul défaut en temps que modèle d'un calcul parallèle est qu'il ne modélise pas de situation de blocage.

Supposons maintenant qu'on décide de n'éliminer que les coupures minimales dans l'ordre: pour éliminer une coupure située après une formule, il faut attendre que celle-ci soit mise en coupure avec sa formule duale. Alors on a modélisé un calcul parallèle avec des situations de blocage.

Nous pensons que cette interprétation pourrait rapidement devenir plus tangible en faisant le lien avec [MTV90].

Notice it is more a strategy than constraints and concurrency, since the calculus always ends, and since our calculus naturally involves a rule ENTROPIE which forgets (part of) the order.

But, when there are extra-logical axioms, one has to leave out this rule. Actually, if these extra-logical axioms express constraints for a robot, such as "open the door before entering", since the order merge with others inside the proof, it would be an obvious blunder to allow anywhere in a proof an application of this ENTROPIE rule.

Thus, with extra logical axioms, this order encodes the temporal dependency and independency of actions, i.e. a concurrent calculus⁴. The only drawback in such a model is that it does not model dead-lock situations.

Now assume we decide to only compute minimal cuts (in the order): to compute a cut above a formula, one has to wait until the formula is itself cut with its dual formula. Then we obtain a model for concurrency with dead locks situations.

We think this interpretation will soon become more concrete by making a connection with [MTV90].

³true concurrency: on ne considère pas que A parallèle à B soit " A avant B ou (non-déterministe) B avant A "; il s'agit plus d'un calcul à la Winskell qu'à la Milner

⁴true concurrency: we do not consider A parallel B to be A before B (non-deterministic) or B before A : it is more like a Winskell calculus than a Milner calculus.

§C. Terminologie et Notations

On trouvera un index des notations, un index et la table des matières à la fin de la thèse.

Pour ce qui est des graphes, on utilise la terminologie du livre de Berge [Ber79] qui donne toutes les définitions dont nous avons besoin dans le premier chapitre, et un index trilingue à la fin de son livre.

On ne s'en écartera que sur le point suivant: un circuit dans un graphe simple comportant et des arcs et des arêtes (des paires d'arcs $\{u = (x, y); \tilde{u} = (y, x)\}$) sera pour nous un circuit (en son sens) c tel que, si $u \in c$ alors $\tilde{u} \notin c$. Cet usage élimine les circuits "plats", comme c'est souvent l'usage (par exemple si on regarde un graphe comme un complexe simplicial de dimension 1).

En s'éloignant de l'usage on dira graphe coloré pour graphe dont les arêtes sont colorées — mais nous ne considérerons jamais de graphes dont les sommets soient colorés, ce qui justifie cet abus commode.

On dira souvent ensemble pour multi-ensemble, et formule pour occurrence de formule, le contexte permettant aisément de trancher.

Ordre signifiera toujours ordre partiel, et relation relation binaire.

La réunion de deux relations u et v de domaines différents E et E' désignera la relation sur $E \cup E'$ dont les couples en relations sont $u \cup v$.

Les intervalles $[n, p]$ sont évidemment pris dans $\mathbb{N}$ et $n[p]$ désigne l'entier n modulo p , ce qui est utile pour indexer les points d'un circuit.

Enfin, tous les objets considérés dans cette thèse sont finis, sauf les espaces cohérents; on remarquera de plus que les résultats démontrés sont rarement vrais dans des structures infinies.

§C. Terminology and Notation

We give a notation index, an index and the table of contents at the end of the thesis.

Concerning graphs we refer to the terminology of [Ber79] which exposes in the first chapter all the definitions we need — and give a multilingual index at the end of his book.

We only depart from his definition of a circuit on the following point: according to our convention a circuit in a simple graph containing both arcs and edges (pairs of arcs $\{u = (x, y); \tilde{u} = (y, x)\}$) is one of his circuit c such that, if $u \in c$ then $\tilde{u} \notin c$. This convention prohibits "flat" circuits as it is usual (e.g. when a graph is viewed as a simplicial complex of dimension 1).

Though it is unusual, we mean by "coloured graph" an edge coloured graph — but we never consider (vertex) coloured graphs, so it is a harmless and convenient abuse.

We often say set for multiset, and formula for occurrence of a formula, the sense being clear from the context.

Order always mean partial order, and relation binary relation.

The union of two relations u and v on two different sets E and E' will be the relation on $E \cup E'$ consisting of the pairs $u \cup v$.

Intervals $[n, p]$ are obviously included in $\mathbb{N}$ and $n[p]$ stand for the integer n modulo p , which is useful for naming the vertices of a circuit.

Finally all objects we consider are finite, except coherence spaces; one can show that most of the results we establish fail for infinite structures.

Partie A

Combinatoire

Chapitre 1

Ordres et Orthogonalité

Comme la conclusion d'un séquent ou d'un réseau ordonné est un multi-ensemble fini ordonné, nous avons préalablement besoin d'opérations élémentaires sur les ordres et d'énoncer quelques propriétés élémentaires de l'orthogonalité des relations binaires, en particulier lorsqu'une des deux est un ordre.

NOTATION 1.1. Par ordre nous entendons ici un ordre dont le domaine est fini. Un tel ordre est totalement déterminé par la donnée, pour tout point x du domaine, de l'ensemble de ses successeurs, noté Sx — ou de l'ensemble de ses prédécesseurs, noté Px .

Par ordre trivial, nous entendons un ordre dont le domaine est un singleton. Comme il n'y a qu'un tel ordre à isomorphisme près, on dira souvent l'ordre trivial.

On écrit $x \leq y[i]$ pour xiy et $x < y[i]$ pour $(xiy \wedge x \neq y)$ ce qui évite d'introduire des notations pour l'ordre strict associé à i .

Par ordre vide, noté $\emptyset$ on entend un ordre discret dont l'ordre strict associé est vide.

On notera $\mathcal{I}\mathcal{D}$ la relation identique, $|u|$ le domaine de la relation u , $\bar{u}$ la clôture réflexive et transitive de u et $u|_E$ la restriction à E de la relation u .

§A. Points Equivalents et Equivalents-Inférieurs

DÉFINITION 1.2. [x équivalent à y , $x \sim y$] On dit que x est équivalent à y dans l'ordre i , ce que l'on note $x \sim y[i]$, si et seulement si

$$\forall z \left[\begin{array}{c} z < x[i] \iff z < y[i] \\ \text{et} \\ x < z[i] \iff y < z[i] \end{array} \right] \quad \text{ce qui s'énonce aussi:} \quad \left[\begin{array}{c} Sx = Sy \\ \text{et} \\ Px = Py \end{array} \right]$$

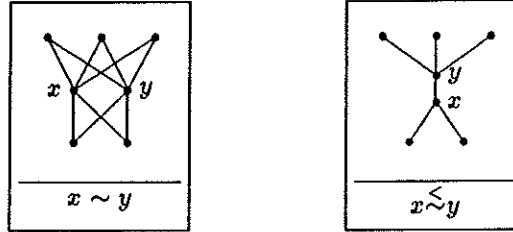
C'est bien évidemment une relation d'équivalence, et si $x \sim y$ alors x et y sont incomparables ou égaux.

DÉFINITION 1.3. [x équivalent-inférieur à y , $x \lesssim y[i]$] On dit que x est équivalent-inférieur à y dans l'ordre i , ce que l'on note $x \lesssim y[i]$, si et seulement si:

$$\forall z \left\{ \begin{array}{l} z \leq x[i] \iff z < y[i] \\ \text{et} \\ x < z[i] \iff y \leq z[i] \end{array} \right\} \quad \text{ce qui s'énonce aussi:} \quad \left\{ \begin{array}{l} Sx = \{y\} \\ \text{et} \\ Py = \{x\} \end{array} \right\}$$

On a alors $x < y$ et cette relation est anti-réflexive, anti-symétrique et anti-transitive. De plus, pour tout x il existe au plus un y tel que $x \lesssim y$.

Voici ce que cela signifie localement dans le diagramme de Hasse:



§B. Contractions d'un Ordre

Nous avons besoin d'opérations sur les ordres qui correspondent aux règles du calcul des séquents, et notamment de définir un ordre sur $A_1, \dots, A_n, A \forall B$ et sur $A_1, \dots, A_n, A < B$, étant donné un ordre i sur $A_1, \dots, A_n, A, B$. On décrit ici ces opérations, qui seront utilisées dans les paragraphes §G..

DÉFINITION 1.4. [contraction d'un ordre suivant deux points équivalents, $i/x \sim y \rightsquigarrow u$]

Soient i une relation d'ordre, x, y deux points de $|i|$ tels que $x \sim y$, et $u \notin |i|$. On définit la contraction de i suivant $x \sim y$, que l'on note $i/x \sim y \rightsquigarrow u$, la relation d'ordre sur $|i| - \{x, y\} \cup \{u\}$ obtenue en identifiant les deux points x et y et en appelant le point résultant u :

$$\begin{array}{ll} u < z[i/x \sim y \rightsquigarrow u] & \iff (x < z[i] \wedge y < z[i]) \\ z < u[i/x \sim y \rightsquigarrow u] & \iff (z < x[i] \wedge z < y[i]) \\ \forall z, z' \neq u \quad z' < z[i/x \sim y \rightsquigarrow u] & \iff z < z'[i] \end{array}$$

On peut aussi formuler cette définition en termes de prédécesseurs:

dans $i/x \sim y \rightsquigarrow u$	dans i
	$Pu = Px = Py$
$\forall z \notin Sx(=Sy)[i]$	$Pz = Pz$
$\forall z \in Sx(=Sy)[i]$	$Pz = Pz - \{x, y\} + \{u\}$

ou de successeurs:

dans $i/x \sim y \rightsquigarrow u$	dans i
$Su = Sx (= Sy)$	
$\forall z \notin Px (= Py)[i]$	$Sz = Sz$
$\forall z \in Px (= Py)[i]$	$Sz = Sz - \{x, y\} + \{u\}$

DÉFINITION 1.5. [contraction d'un ordre suivant deux points équivalents-inférieurs, $i/u \rightsquigarrow x < y$]

Soient i une relation d'ordre, x, y deux points de $|i|$ tels que $x \lesssim y$, et $u \notin |i|$. On définit la contraction d'un ordre suivant deux points équivalents-inférieurs, que l'on note $i/x \lesssim y \rightsquigarrow u$ la relation d'ordre sur $|i| - \{x, y\} \cup \{u\}$ obtenue en identifiant les deux points x et y et en appelant le point résultant u :

$u < z[i/x \lesssim y \rightsquigarrow u]$	$\iff (x < z[i] \wedge y < z[i])$
$z < u[i/x \lesssim y \rightsquigarrow u]$	$\iff (z < x[i] \wedge z < y[i])$
$\forall z, z' \neq u \quad z' < z[i/x \lesssim y \rightsquigarrow u]$	$\iff z < z'[i]$

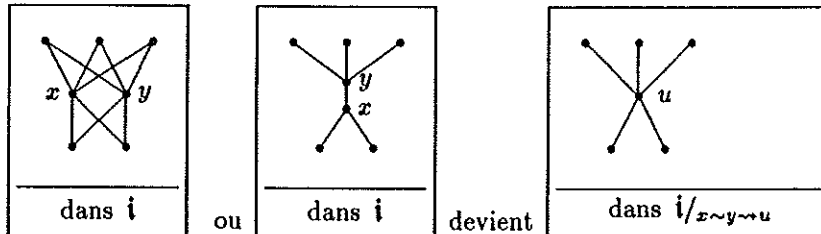
On peut aussi formuler cette définition en termes de prédécesseurs:

dans $i/x \lesssim y \rightsquigarrow u$	dans i
$Pu = Px$	
$\forall z \notin Sy[i]$	$Pz = Pz$
$\forall z \in Sy[i]$	$Pz = Pz - \{y\} + \{u\}$

ou de successeurs:

dans $i/x \lesssim y \rightsquigarrow u$	dans i
$Su = Sy$	
$\forall z \notin Px[i]$	$Sz = Sz$
$\forall z \in Px[i]$	$Sz = Sz - \{x\} + \{u\}$

Résumons graphiquement ces deux opérations:



REMARQUE 1.6. (i) Les ordres induits sur $|i| - \{x, y\}$ par i , $i/x \sim y \rightsquigarrow u$ et $i/x \lesssim y \rightsquigarrow u$ coïncident.

(ii) On a z équivalent à u dans $i/x \sim y \rightsquigarrow u$ si et seulement si x, y et z sont équivalents dans i . De plus les ordres $(i/x \sim y \rightsquigarrow u) /_{u \sim z \rightsquigarrow w}$ et $(i/x \sim z \rightsquigarrow u) /_{u \sim y \rightsquigarrow w}$ sont les mêmes.

(iii) Soit x équivalent-inférieur à y dans i . On a $u \lesssim z$ dans $i/x \lesssim y \rightsquigarrow u$ si et seulement si x est équivalent-inférieur à y dans i et y est équivalent inférieur à z dans i . De plus les ordres

$(i/x \sim y \rightsquigarrow u) /_{u \sim z \rightsquigarrow w}$ et $(i/y \sim z \rightsquigarrow v) /_{x \sim u \rightsquigarrow w}$ sont les mêmes.

(iv) Si x, y, z et t sont quatre points distincts de i , tels que les couples x, y et z, t soient chacun un couple de points équivalents ou équivalents-inférieurs de i , alors les ordres $(i/x, y \rightsquigarrow u) /_{z, t \rightsquigarrow v}$ et $(i/z, t \rightsquigarrow v) /_{x, y \rightsquigarrow u}$ sont les mêmes.

(v) Si x et y sont deux points équivalents-inférieurs, et z et t sont deux points distincts équivalents de i , alors les quatre points x, y, z et t de i sont distincts.

Soit E un ensemble; appelons $(E, \mathfrak{V}, <)$ les formules écrites à l'aide des éléments de E et des connecteurs binaires $\mathfrak{V}$ et $<$. On quotiente $(E, \mathfrak{V}, <)$ par l'associativité des deux connecteurs $\mathfrak{V}$ et $<$, et par la commutativité de $\mathfrak{V}$ (ces connecteurs coïncideront bien évidemment avec ceux de même nom de la troisième partie, mais ce qui est présenté ici, peut se concevoir de manière purement formelle).

Soit i un ordre sur E ; si $x \sim y$, on note $(x \mathfrak{V} y)$ le point résultant de la contraction de x et y dans $i/x \rightsquigarrow y \rightsquigarrow u$, et si $x < y$, on note $(x < y)$ le point résultant de la contraction de x et y dans $i/x \rightsquigarrow y \rightsquigarrow u$. Remarquons que les domaines des contractions successives de i sont ainsi constitués de formules de $(E, \mathfrak{V}, <)$ où chaque élément de E apparaît exactement une fois dans exactement une de ces formules.

Soient i_1 et i_2 deux ordres sur des formules de $(E, \mathfrak{V}, <)$. On dit que i_1 et i_2 sont équivalents s'il existe une bijection du domaine de i_1 dans le domaine de i_2 qui préserve l'ordre, et telle que l'image d'une formule soit une formule équivalente.

La proposition suivante découle des remarques 1.6., et établit la confluence locale de la contraction:

PROPOSITION 1.7. Soient i_1 et i_2 deux ordres équivalents dont les domaines, constitués de formules de $(E, \mathfrak{V}, <)$ sont tels que chaque élément de E apparaisse exactement une fois dans exactement une formule; soit j_1 une contraction de i_1 et j_2 une contraction de i_2 . Si j_1 n'est pas équivalent à j_2 , alors il existe k_1 une contraction de j_1 et k_2 une contraction de j_2 qui sont équivalentes.

Un argument standard permet alors d'en déduire le

COROLLAIRE 1.8. Soient i et j deux ordres sur E et i_0 et j_0 deux ordres obtenus par des contractions successives de i et j tels que i_0 et j_0 ne contiennent ni points équivalents (distincts) ni équivalents-inférieurs. Alors

$$i = j \iff i_0 \text{ et } j_0 \text{ sont équivalents}$$

§C. Expansions d'un Ordre

Etant donné un ordre sur des formules et des coupures, nous devons lui associer, lors de la description des étapes élémentaires d'élimination des coupures, (§B.1. et §B. du chapitre 7) un ordre dont le domaine s'obtient en remplaçant une des coupures par deux coupures plus petites. Nous avons besoin de deux opérations distinctes, suivant qu'il s'agit d'une coupure entre deux

précède ou entre un par et un tenseur . On remarquera que les opérations dont nous donnons ici la description sont les opérations inverses des contractions précédentes.

DÉFINITION 1.9. [expansion équivalente d'un ordre, $i/u \rightsquigarrow x \sim y$] Soient i une relation d'ordre, $u \in |i|$ et $x, y \notin |i|$. On définit l'expansion équivalente de i suivant u , que l'on note $i/u \rightsquigarrow x \sim y$, comme la relation d'ordre sur $|i| - \{u\} \cup \{x, y\}$ obtenue en remplaçant u par deux nouveaux points équivalents x et y :

$$\begin{array}{l} x \sim y [i/u \rightsquigarrow x \sim y] \\ x < z [i/u \rightsquigarrow x \sim y] \iff (u < z [i]) \\ z < x [i/u \rightsquigarrow x \sim y] \iff (z < u [i]) \\ \forall z, z' \neq x, y \quad z' < z [i/u \rightsquigarrow x \sim y] \iff z < z' [i] \end{array}$$

ce qui peut aussi se formuler en termes de prédécesseurs:

$$\begin{array}{ll} \text{dans } i/u \rightsquigarrow x \sim y & \text{dans } i \\ Px = Py & = Pu \\ \forall z \notin Su & Pz = Pz \\ \forall z \in Su & Pz = Pz - \{u\} + \{x, y\} \end{array}$$

ou de successeurs:

$$\begin{array}{ll} \text{dans } i/u \rightsquigarrow x \sim y & \text{dans } i \\ Sx = Sy & = Su \\ \forall z \notin Pu & Sz = Sz \\ \forall z \in Pu & Sz = Sz - \{u\} + \{x, y\} \end{array}$$

DÉFINITION 1.10. [expansion équivalente-inférieure d'un ordre, $i/u \rightsquigarrow x < y$] Soient i une relation d'ordre, $u \in |i|$ et $x, y \notin |i|$. On définit l'expansion équivalente de i suivant u , que l'on note $i/u \rightsquigarrow x < y$, comme la relation d'ordre sur $|i| - \{u\} \cup \{x, y\}$ obtenue en remplaçant u par deux nouveaux points équivalents x et y :

$$\begin{array}{l} x \lessdot y [i/u \rightsquigarrow x < y] \\ y < z [i/u \rightsquigarrow x < y] \iff (u < z [i]) \\ z < x [i/u \rightsquigarrow x < y] \iff (z < u [i]) \\ \forall z, z' \neq x, y \quad z' < z [i/u \rightsquigarrow x < y] \iff z < z' [i] \end{array}$$

ce qui peut aussi se formuler en termes de prédécesseurs:

$$\begin{array}{ll} \text{dans } i/u \rightsquigarrow x < y & \text{dans } i \\ Px & = Pu \\ Py & = \{x\} \\ \forall z \notin Su & Pz = Pz \\ \forall z \in Su & Pz = Pz - \{u\} + \{y\} \end{array}$$

ou de successeurs:

	dans $i/u \rightsquigarrow x \rightsquigarrow y$	dans i
	$Sx = \{y\}$	
	$Sy = Su$	
$\forall z \notin Pu$	$Sz = Sz$	
$\forall z \in Pu$	$Sz = Sz - \{u\} + \{x\}$	

Ces opérations d'expansion sont "inverses" des précédentes opérations de contraction:

$$\begin{aligned}
 (i/u \rightsquigarrow x \rightsquigarrow y) /_{x \rightsquigarrow y \rightsquigarrow u} &= i \\
 (i/u \rightsquigarrow x < y) /_{x \rightsquigarrow y \rightsquigarrow u} &= i \\
 (i/x \rightsquigarrow y \rightsquigarrow u) /_{u \rightsquigarrow x \rightsquigarrow y} &= i \\
 (i/x \rightsquigarrow y \rightsquigarrow u) /_{u \rightsquigarrow x \rightsquigarrow y} &= i
 \end{aligned}$$

§D. Ordres Contractiles

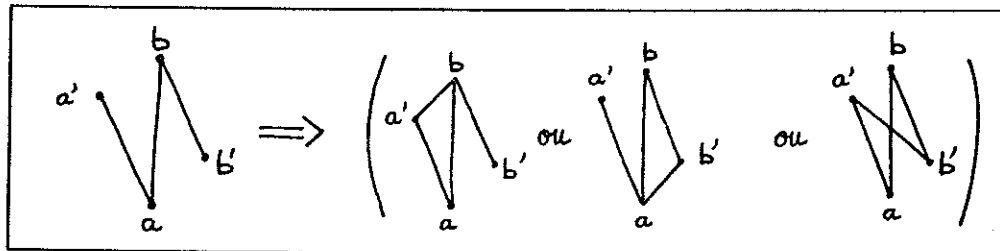
DÉFINITION 1.11. On dit qu'un ordre est contractile si et seulement si il se réduit en l'ordre trivial par une suite (finie) de contractions.

REMARQUE 1.12. Dès que le domaine a plus de trois points, il existe des ordres non contractiles: $i = \{a < b, b' < b, a < a'\}$ est le plus petit ordre qui n'est pas contractile. Le théorème suivant montre que ce contre-exemple est tout à fait générique.

La caractérisation des ordres contractiles est la suivante:

THÉORÈME 1.13. Un ordre fini i est contractile si et seulement s'il satisfait la condition globale $\mathbb{P}^1$ suivante:

$$\mathbb{P}^1 : \quad \forall a, b, a', b' \left(\begin{array}{l} a < a' \\ \text{et} \\ a < b \\ \text{et} \\ b' < b \end{array} \right) \Rightarrow \left(\begin{array}{l} a' \leq b \\ \text{ou} \\ a \leq b' \\ \text{ou} \\ b' \leq a' \end{array} \right)$$



¹comme "papillon": le cas utile dans la conclusion est $b' \leq a'$ et dans le diagramme de Hasse, le dessin obtenu ressemble à un papillon

REMARQUE 1.14.

- (i) L'implication est vraie dès que deux des quatre points sont égaux.
- (ii) $\mathbf{i}$ satisfait cette condition, si et seulement si, pour toute partie E de $|\mathbf{i}|$, l'ordre induit par $\mathbf{i}$ sur E satisfait cette condition.
- (iii) Un ordre satisfait la condition $\mathbb{P}$ si et seulement si l'ordre inverse la satisfait.
- (iv) L'exemple $\mathbf{i} = \{(x_i < x_{i+1}), (x_i < y_i), i \in \mathbb{N}\}$ montre que l'hypothèse de finitude est essentielle: $\mathbf{i}$ satisfait $\mathbb{P}$ mais n'est pas contractile.

Commençons par quelques lemmes:

LEMME 1.15. Les opérations de contraction et d'expansion préservent $\mathbb{P}$. Autrement dit, si $\mathfrak{k}$ est une contraction de $\mathbf{i}$ alors $\mathfrak{k}$ satisfait $\mathbb{P}$ si et seulement si $\mathbf{i}$ satisfait $\mathbb{P}$.

DÉMONSTRATION:

" $\mathbf{i}$ satisfait $\mathbb{P}$ " $\Rightarrow$ " $\mathfrak{k}$ satisfait $\mathbb{P}$ "

Soient a, b, a', b' des points deux à deux distincts de $\mathfrak{k}$ satisfaisant $a < b \wedge a < a' \wedge b' < b$ dans $\mathfrak{k}$, montrons que $a \leq b' \vee a' \leq b \vee b' \leq a'$ dans $\mathfrak{k}$.

Si aucun des a, b, a', b' n'est le point u obtenu par contraction de x et y , le point (ii) de la remarque précédente permet de conclure en considérant les restrictions de $\mathbf{i}$ et $\mathfrak{k}$ à $\{a, b, a', b'\}$ qui coïncident.

Supposons donc que l'un des a, b, a', b' soit u . Par symétrie (en changeant le sens de l'ordre) on voit qu'il n'y a que deux cas à envisager: $a = u$ et $a' = u$.

$a = u$ On a donc:

$$u < b[\mathfrak{k}] \text{ et } a < a'[\mathfrak{k}] \text{ et } b' < b[\mathfrak{k}]$$

C'est donc qu'on avait, dans $\mathbf{i}$:

$$y < b[\mathbf{i}] \text{ et } a < a'[\mathbf{i}] \text{ et } b' < b[\mathbf{i}].$$

En appliquant la condition $\mathbb{P}$ dans $\mathbf{i}$, il vient:

$$y \leq b'[\mathbf{i}] \text{ ou } a' \leq b[\mathbf{i}] \text{ ou } b' \leq a'[\mathbf{i}]$$

ce qui entraîne, dans $\mathfrak{k}$

$$y \leq b'[\mathfrak{k}] \text{ ou } a' \leq b[\mathfrak{k}] \text{ ou } b' \leq a'[\mathfrak{k}]$$

$a' = u$ Tout aussi simple que le cas précédent.

" $\mathfrak{k}$ satisfait $\mathbb{P}$ " $\Rightarrow$ " $\mathbf{i}$ satisfait $\mathbb{P}$ " Soient a, b, a', b' des points deux à deux distincts de $\mathbf{i}$ satisfaisant $a < b \wedge a < a' \wedge b' < b$ dans $\mathbf{i}$, montrons que $a \leq b' \vee a' \leq b \vee b' \leq a'$ dans $\mathbf{i}$. Si aucun de ces quatre points n'est l'un des deux points contractés, comme la contraction préserve l'ordre sur les points autres que ceux contractés, et que dans $\mathfrak{k}$ l'implication est vraie, elle l'est aussi dans $\mathbf{i}$.

Si $\mathfrak{k}$ est obtenu par contraction de deux points équivalents-inférieurs $x \lesssim y$ de $\mathbf{i}$ il est possible que deux de ces quatre points soient les deux points contractés. Par symétrie, il suffit d'envisager les deux cas suivants: $a = x$ et $b = y$, et $a = x$ et $a' = y$.

Si $a = x$ et $b = y$, comme ces points sont équivalents-inférieurs, on a dans i : $b' \leq x$ et $a' \geq y$, d'où $b' < a'$ dans i .

Si $a = x$ et $a' = y$, comme ces points sont équivalents-inférieurs, on a dans i : $b \geq y$.

On peut donc désormais supposer qu'un et un seul des quatre points a, b, a', b' est x ou y et se restreindre à ne considérer que les deux cas suivants:

si $a = x$ ou y on a $u < b \wedge a < a' \wedge b' < b$ dans $\mathbb{E}$, et il s'ensuit que $u \leq b' \vee a' \leq b \vee b' \leq a'$ dans $\mathbb{E}$; on a alors $(x \leq b' \wedge y \leq b') \vee a' \leq b \vee b' \leq a'$ dans i , ce qui assure le résultat.

si $a' = x$ ou y on a $a < b \wedge a < u \wedge b' < b$ dans $\mathbb{E}$, et il s'ensuit que $a \leq b' \vee u \leq b \vee b' \leq a'$ dans $\mathbb{E}$; on a alors $a \leq b' \vee a' \leq b \vee (b' \leq x \wedge b' \leq y)$ dans i , ce qui assure le résultat.

◇

LEMME 1.16. Soient i satisfaisant $\mathbb{P}$, $x_1 < \dots < x_m$ une chaîne de longueur maximale $m \geq 2$, y_1 minorant strict de x_2 et y_m un majorant strict de x_{m-1} . Alors $x_1 \sim y_1$ et $x_m \sim y_m$.

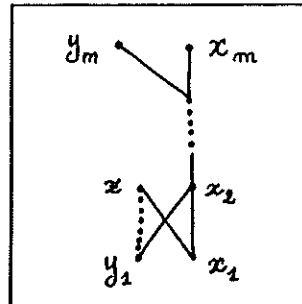
DÉMONSTRATION: Comme l'ordre inverse satisfait aussi $\mathbb{P}$ il suffit de montrer que $x_1 \sim y_1$.

Remarquons pour commencer que y_1 est un prédécesseur de x_2 : s'il existait un w tel que $y_1 < w < x_2$ la chaîne $y_1 < \dots < w < \dots < x_2 < \dots < x_m$ serait de longueur strictement supérieure à m .

- * Comme la chaîne $x_1 < x_2 < \dots < x_m$ est de longueur maximale, x_1 n'a pas de minorant strict; la chaîne $y_1 < x_2 < \dots < x_m$ étant aussi de longueur maximale y_1 n'a pas non plus de minorant strict.
- * Soit $z > x_1$, montrons que $z > y_1$. On a $z > x_1$, $x_1 < x_2$ et $y_1 < x_2$, et la condition $\mathbb{P}$ montre que

$$x_1 \leq y_1 \vee z \leq x_2 \vee y_1 \leq z.$$

Les deux premiers cas sont exclus car x_2 est un successeur de x_1 . C'est donc que $y_1 \leq z$, et comme x_2 est un successeur de x_1 , $y_1 \neq z$, et on a établi $y_1 < z$. La chaîne $y_1 < x_2 < \dots < x_m$ étant aussi de de longueur maximale, on montre de même que, si $z > y_1$, alors $z > x_1$.



Ainsi x_1 et y_1 ont les mêmes minorants et majorants stricts, i.e. $x_1 \sim y_1$.

◇

LEMME 1.17. Soit $\mathfrak{i}$ un ordre satisfaisant $\mathbb{P}$. Si $Px \cap Py \neq \emptyset$ et $Sx \cap Sy \neq \emptyset$ alors $x \sim y$.

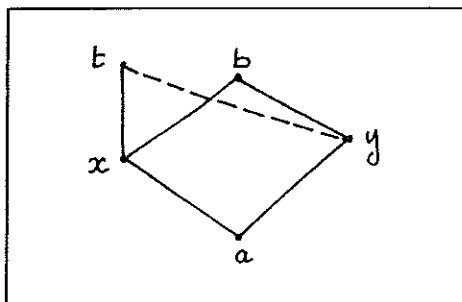
DÉMONSTRATION: Soient $a \in Px \cap Py$ et $b \in Sx \cap Sy$ et soit t un successeur de x . Comme on a

$$x < t \wedge x < b \wedge y < b$$

la condition $\mathbb{P}$ impose que:

$$t \leq b \vee x \leq y \vee y \leq t$$

On voit que $t \leq b$ contredit le fait que t soit un successeur de x , et que $x \leq y$ contredit le fait que b soit un successeur de y . Ainsi $y \leq t$.



Les rôles de x et de y étant symétriques, on montrerait de même que tout successeur de y est supérieur ou égal à x . Ainsi x et y ont les mêmes majorants stricts.

La même démonstration appliquée à l'ordre inverse prouve que x et y ont aussi les mêmes minorants stricts. Ils sont donc équivalents. $\diamond$

Démonstration du théorème 1.13.: Montrons pour commencer que la condition $\mathbb{P}$ est nécessaire, i.e. si $\mathfrak{i}$ est contractile, alors $\mathfrak{i}$ satisfait $\mathbb{P}$. L'ordre trivial satisfaisant $\mathbb{P}$, ceci est une conséquence directe du lemme 1.15..

Etablissons maintenant que si $\mathfrak{i}$ satisfait $\mathbb{P}$, alors $\mathfrak{i}$ est contractile. On procède par récurrence sur le nombre de points du domaine de $\mathfrak{i}$: il suffit alors, en vertu du lemme 1.15., de montrer que tout ordre non trivial et satisfaisant $\mathbb{P}$ contient deux points distincts équivalents ou équivalents-inférieurs. Soit $x_1 < x_2 < \dots < x_m$ une chaîne de longueur maximale m .

* Si $m = 1$, l'ordre $\mathfrak{i}$ est l'ordre discret (qui satisfait $\mathbb{P}$) dont tous les points sont équivalents; comme $\mathfrak{i}$ n'est pas l'ordre trivial, $\mathfrak{i}$ contient deux points équivalents.

* Si $m \geq 2$,

▷ OU BIEN POUR TOUT $k \leq m - 1$ LE POINT x_k A UN SUCCESSEUR $y_{k+1} \neq x_{k+1}$, et dans ce cas le lemme 1.16. montre que $x_m \sim y_m$.

▷ OU BIEN IL EXISTE DES x_k AVEC $k \leq m - 1$ DONT LE SEUL SUCCESSEUR EST x_{k+1} , ET SOIT x_i LE PLUS PETIT D'ENTRE EUX.

Si $i = 1$,

- ou bien x_1 est l'unique prédécesseur de x_2 et $x_1 \lesssim x_2$.

- ou bien x_2 a un prédécesseur $y_1 \neq x_1$ et le lemme 1.16. montre que $x_1 \sim y_1$.

Si $i > 1$, le point x_{i-1} a un successeur autre que x_i disons v .

Si x_{i+1} n'a pas d'autre prédécesseur que x_i ,

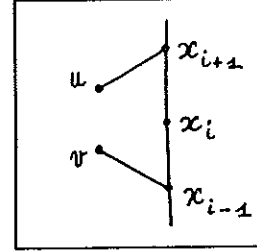
on a alors $x_i \lesssim x_{i+1}$.

Sinon, soit $u \neq x_i$ un prédécesseur de x_{i+1} On a alors

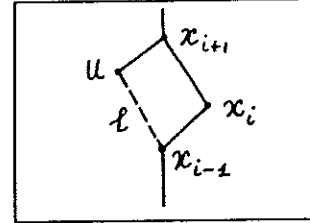
$$u < x_{i+1} \wedge x_{i-1} < x_{i+1} \wedge x_{i-1} < v$$

et la condition $\mathbb{P}$ montre que:

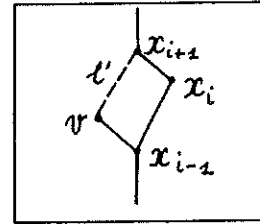
$$x_{i-1} \leq u \vee v \leq x_{i+1} \vee u \leq v.$$



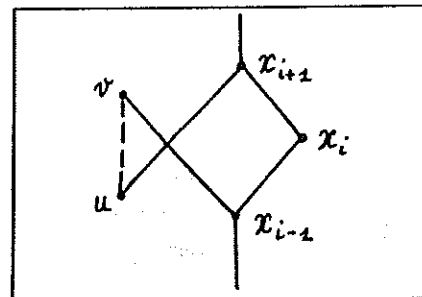
Si $x_{i-1} \leq u$, appelons l la longueur de la chaîne c de x_{i-1} à u . On a alors $l \leq 1$. En effet, si $l > 1$ la chaîne $x_1 < x_2 \dots < x_{i-1} < x_i < x_{i+1} < \dots < x_m$ ne peut être maximale car la chaîne $x_1 < x_2 \dots < x_{i-1} < c < u < x_{i+1} < \dots < x_m$ est strictement plus longue. Donc $l = 1$, $x_{i+1} \in Sx_i \cap Su$ et $x_{i-1} \in Px_i \cap Pu$, et le lemme 1.17. montre que $u \sim x_i$.



Si $v \leq x_{i+1}$, appelons l' la longueur de la chaîne c' de v à x_{i+1} . On a alors $l' \leq 1$. En effet, si $l' > 1$ la chaîne $x_1 < x_2 \dots < x_{i-1} < x_i < x_{i+1} < \dots < x_m$ ne peut être maximale car la chaîne $x_1 < x_2 \dots < x_{i-1} < v < c' < x_{i+1} < \dots < x_m$ est strictement plus longue. Donc $l' = 1$, $x_{i+1} \in Sx_i \cap Sv$ et $x_{i-1} \in Px_i \cap Pv$, et le lemme 1.17. montre que $x_i \sim v$.



Si $u \leq v$ et $\neg(x_i \leq v)$, montrons que $u = v$ i.e. que $\neg(u < v)$. Si on avait $u < v$, comme $x_i < x_{i+1} \wedge u < x_{i+1}$, on aurait, en vertu de la condition $\mathbb{P}$, $v \leq x_{i+1} \vee u \leq x_i \vee x_i \leq v$. Cependant, comme u est un prédécesseur de x_{i+1} on a $\neg(u \leq x_i)$, comme v est un successeur de x_{i-1} on a $\neg(v \leq x_i)$, et nous avons exclu le cas (déjà traité) $v \leq x_{i+1}$. On a donc $u = v$ et le lemme 1.17. montre que $(u = v) \sim x_i$.



Ce résultat nous permettra de montrer dans le paragraphe §G. du chapitre 7 que les connecteurs généralisés définis par un ordre définissable à l'aide des connecteurs binaires sont exactement ceux définis par un ordre contractile.

On peut maintenant appliquer le résultat 1.8. aux ordres contractiles:

THÉORÈME 1.18. *Un ordre contractile est complètement caractérisé par une formule écrite avec les points de son domaine et les connecteurs $\mathfrak{R}$ et $<$, à la commutativité de $\mathfrak{R}$ et à l'associativité de $\mathfrak{R}$ et $<$ près.*

Ce résultat nous permettra dans le paragraphe §G. du chapitre 7 de montrer que les connecteurs généralisés définissables par des connecteurs binaires le sont d'une unique manière.

§E. Orthogonalité des Relations Binaires

NOTATION 1.19. Soit E un ensemble fini. On note $\mathcal{R}_E$ l'ensemble des relations binaires sur E , $\mathcal{I}_E$ l'identité sur E , et l'on dira simplement relation pour relation binaire. Les relations sont désignées par des lettres gothiques $u, v, \dots$ etc. On note uv la composée des relations binaires u et v (i.e. $x(uv)y$ ssi $\exists z [xuz \wedge zvy]$) et $\bar{u}$ la clôture réflexive et transitive de u : $\bar{u} = \bigcup_{n \in \mathbb{N}} u^n$. Par chaîne d'une relation u , on entend une suite de points (x_i) , de cardinal fini, telle que $x_i u x_{i+1}$. Par circuit d'une relation u , on entend une chaîne de u de cardinal ≥ 2 , dont le premier et le dernier point sont égaux.

DÉFINITION 1.20. On définit sur $\mathcal{R}(E)$ une relation binaire notée " $\perp$ " appelée orthogonalité faible par:

$$u \perp v \text{ ssi } \exists n (uv)^n = \emptyset$$

Soit $S \subset \mathcal{R}_E$ on note

$$S^\perp = \{u \in \mathcal{R}_E / \forall v \in S \ u \perp v\} \subset \mathcal{R}_E$$

et si $S, T \subset \mathcal{R}_E$ on dit que $S \perp T$ ssi $\forall u \in S \ \forall v \in T \ u \perp v$, i.e. $S \subset T^\perp$ ou $T \subset S^\perp$

PROPOSITION 1.21. Deux relations u et v sur E sont faiblement orthogonales si et seulement si leur composée uv ne contient pas de circuit.

DÉMONSTRATION: Soit w une relation sur E . Pour établir la proposition il suffit de montrer que pour toute relation w sur E , on a:

$$(\forall n \ w^n \neq \emptyset) \iff (w \text{ admet un circuit})$$

Si w admet un circuit $x_0, x_1, \dots, x_p, x_0$, alors pour tout entier n , on a $x_0 w x_{n[p]}$, et par suite w^n n'est pas vide.

Inversement, si n est un entier supérieur ou égal à $\#E$ pour lequel $w^n \neq \emptyset$, alors, par définition de la composition des relations, on peut trouver une chaîne de longueur $n+1$ $x_0, x_1, \dots, x_n$; les x_i ne pouvant être deux à deux distincts, cette chaîne contient un circuit de w . $\diamond$

Enumérons brièvement quelques propriétés de l'orthogonalité faible:

REMARQUE 1.22. L'orthogonalité faible jouit des propriétés suivantes, qui justifient le nom "orthogonalité":

★ La relation $\perp$ est symétrique, tant sur $\mathcal{R}_E$ que sur $\mathcal{P}(\mathcal{R}_E)$:

$$(u \vee v)^n = \emptyset \Rightarrow (v \vee u)^{n+1} = v(u \vee v)^n u = \emptyset$$

★ Soit $S, T \subset \mathcal{R}_E$.

$$S \subset T \Rightarrow T^\perp \subset S^\perp$$

$$S \perp T \iff S \subset T^\perp \iff T \subset S^\perp$$

$$(S^\perp)^\perp \supset S$$

$$\left((S^\perp)^\perp \right)^\perp = S^\perp$$

Nous présentons maintenant une autre notion d'orthogonalité qui est essentiellement semblable à la précédente pour les préordres, et qui est stable par clôture réflexive et transitive.

DÉFINITION 1.23. Soient u et v deux relations sur E . On dit que u est orthogonale à v , ce qu'on note $u \perp v$, si et seulement si:

$$(\bar{u} - \mathcal{I}\mathcal{D}) \perp (\bar{v} - \mathcal{I}\mathcal{D})$$

PROPOSITION 1.24. Soient u et v deux relations sur E . Les relations u et v sont orthogonales si et seulement si tout circuit élémentaire de $u \cup v - \mathcal{I}\mathcal{D}$ est un circuit de $u - v$ ou de $v - u$.

DÉMONSTRATION: Soit $x_0, \dots, x_{n-1}, x_0$ un circuit élémentaire de $u \cup v - \mathcal{I}\mathcal{D}$ qui n'est ni un circuit de $u - v$, ni un circuit de $v - u$. Notons que les x_i sont deux à deux distincts et que, quitte à effectuer une permutation circulaire des indices, on peut supposer que $x_0 u x_1$ et $x_{n-1} v x_0$. On peut alors extraire une sous-suite $x_{i_0} = x_0, x_{i_2}, \dots, x_{i_p}$ de $x_0, \dots, x_{n-1}, x_0$ satisfaisant

$$x_{i_0}(\bar{u} - \mathcal{I}\mathcal{D})x_{i_1}, x_{i_1}(\bar{v} - \mathcal{I}\mathcal{D})x_{i_2}, x_{i_{p-1}}(\bar{u} - \mathcal{I}\mathcal{D})x_{i_p}, x_{i_p}(\bar{v} - \mathcal{I}\mathcal{D})x_{i_0}$$

et cela montre que u et v ne sont pas orthogonales.

Inversement, supposons que $(\bar{u} - \mathcal{I}\mathcal{D})(\bar{v} - \mathcal{I}\mathcal{D})$ contienne un circuit. Il existe alors une suite $x_0, \dots, x_{n-1}$ de points de E et une suite $i_0 = x_0 < i_2 \dots < i_p$ d'entiers $\leq n-1$ vérifiant les conditions suivantes, où I_k désigne l'intervalle $[i_k, i_{k+1}]$ si $k < p$ et I_p l'ensemble $\{i_p, i_{p+1}, \dots, i_{n-1}, i_0\}$:

- ★ deux points de la suite $x_0, \dots, x_{n-1}$ dont les indices sont dans un même I_k sont distincts.
- ★ si k est pair et $l, l+1 \in I_k$ alors $x_l u x_{l+1}$
- ★ si k est impair et $l, l+1 \in I_k$ alors $x_l v x_{l+1}$

Si les points $(x_i)_{i \leq n-1}$ sont deux à deux distincts, alors $x_0, \dots, x_{n-1}, x_0$ est un circuit minimal de $u \cup v - \mathcal{I}\mathcal{D}$ qui n'est ni un circuit de $u - v$, ni un circuit de $v - u$, sinon il existe un plus petit indice i et un indice j tels que $i < j < n-1$ et $x_i = x_j$. Comme i et j ne peuvent être

dans le même I_k , $x_i, x_{i+1}, \dots, x_j$ est un circuit élémentaire qui n'est ni un circuit de $u - v$, ni un circuit de $v - u$. $\diamond$

REMARQUE 1.25. $\star u$ et v sont orthogonales si et seulement si leurs clôtures réflexives et transitives $\bar{u}$ et $\bar{v}$ le sont.

$\star u$ et v peuvent être orthogonales et contenir l'une et l'autre des circuits et être orthogonales.

$\star$ si $u \cap v \cap \mathcal{D} = \emptyset$ et si u et v sont orthogonales alors u et v sont faiblement orthogonales.

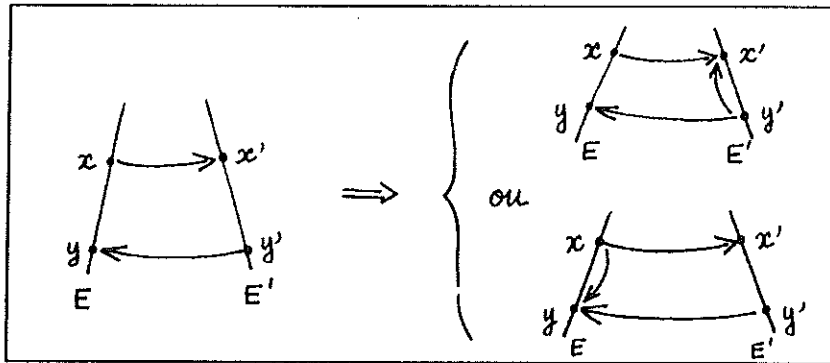
$\star$ l'orthogonalité jouit également des propriétés mentionnées dans la remarque 1.22.

§F. Ordres et Orthogonalité

Les résultats présentés ici servent principalement à définir l'ordre maximal que l'on peut prendre sur le séquent conclusion des règles TENSEUR et MÉLANGE du chapitre 8 — i.e. le plus grand ordre qui ne crée pas de circuit dans les réseaux ordonnés correspondants et qui soit interprété par la sémantique cohérente.

DÉFINITION 1.26. Soient E et E' deux ensembles finis disjoints. On dit qu'une relation $\vdash$ sur $E \cup E'$ satisfait la condition $\mathbb{T}(\vdash, E, E')$ lorsque:

$$\forall x, y \in E \quad \forall x', y' \in E' \quad \left| \begin{array}{c} x \vdash x' \\ \text{et} \\ y' \vdash y \end{array} \right| \Rightarrow \left| \begin{array}{c} x \vdash y \\ \text{ou} \\ y' \vdash x' \end{array} \right|$$



THÉORÈME 1.27. Soient E et E' deux ensembles finis disjoints, $\vdash$ un ordre strict sur E et $\vdash'$ un ordre strict sur E' , et $\vdash$ un ordre strict sur $E \cup E'$ contenant $\vdash \cup \vdash'$. Les propositions suivantes sont équivalentes:

(i) $\forall u \in \mathcal{R}(E) \quad \forall u' \in \mathcal{R}(E') \quad (u \perp \vdash \text{ et } u' \perp \vdash') \Rightarrow \vdash \perp (u \cup u')$

(ii) $\mathbb{T}(\vdash, E, E')$ et $\vdash|_E = \vdash$ et $\vdash|_{E'} = \vdash'$

Si elles sont satisfaites on dit que $\vdash$ est une extension orthogonale de $\vdash, \vdash'$

DÉMONSTRATION: Montrons que (i) implique (ii).

Pour établir que les restrictions de $\mathbf{l}$ à E et E' sont $\mathbf{i}$ et $\mathbf{i}'$, compte tenu de l'inclusion $\mathbf{i} \cup \mathbf{i}' \subset \mathbf{l}$ et de la symétrie du problème, il suffit de montrer que $\mathbf{i}$ contient la restriction de $\mathbf{l}$ à E . Si ce n'était pas le cas, il existerait de points distincts x et y de E tels que $x \mathbf{l} y$ sans que $x \mathbf{i} y$. La relation $\mathbf{u} = \{(y, x)\} = \mathbf{u} \cup \emptyset$ serait orthogonale à $\mathbf{i}$ mais non à $\mathbf{l}$.

Supposons maintenant que $\mathbf{T}(\mathbf{l}, E, E')$ ne soit pas satisfaite. Il existe donc $x, y \in E$ et $x', y' \in E'$ tels que $x \mathbf{l} x', y' \mathbf{l} y, \neg(x \mathbf{l} y)$ et $\neg(y' \mathbf{l} x')$. Notons que $x \neq y$ et $x' \neq y'$ puisque $\mathbf{l}$ est transitive. Comme les restrictions de $\mathbf{l}$ à E et E' sont respectivement $\mathbf{i}$ et $\mathbf{i}'$, il s'ensuit que $\mathbf{u} = \{(y, x)\}$ et $\mathbf{u}' = \{(x', y')\}$ sont respectivement orthogonales à $\mathbf{i}$ et $\mathbf{i}'$. Or x, x', y', y, x est un circuit élémentaire de $\mathbf{l} \cup (\mathbf{u} \cup \mathbf{u}') - \mathcal{I} \emptyset$ qui n'est ni un circuit de $\mathbf{l} - (\mathbf{u} \cup \mathbf{u}')$, ni un circuit de $(\mathbf{u} \cup \mathbf{u}') - \mathbf{l}$, ce qui montre que $\mathbf{l}$ et $\mathbf{u} \cup \mathbf{u}'$ ne sont pas orthogonales.

Montrons maintenant que (ii) implique (i). Supposons qu'il existe deux relations $\mathbf{u}$ et $\mathbf{u}'$ qui soient respectivement orthogonales à $\mathbf{i}$ et $\mathbf{i}'$ et telles que $\mathbf{u} \cup \mathbf{u}'$ ne soit pas orthogonale à $\mathbf{l}$. Il existe alors des circuits de $(\mathbf{u} \cup \mathbf{u}') \cup \mathbf{l} - \mathcal{I} \emptyset$ élémentaires qui ne sont ni des circuits de $(\mathbf{u} \cup \mathbf{u}') - \mathbf{l}$, ni des circuits de $\mathbf{l} - (\mathbf{u} \cup \mathbf{u}')$. Choisissons en un qui soit de longueur minimale.

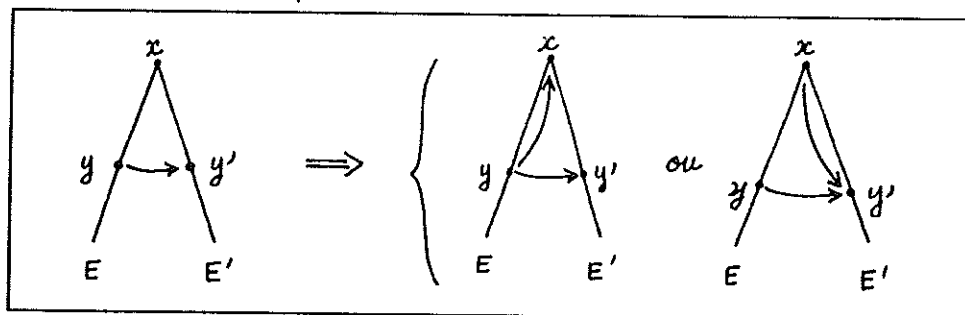
* S'il ne contient que des points de E , c'est un circuit de $(\mathbf{u} \cup \mathbf{l}|_E - \mathcal{I} \emptyset)$ qui n'est ni un circuit de $\mathbf{u} - \mathbf{l}|_E$, ni un circuit de $\mathbf{l}|_E - \mathbf{u}$, et par suite $\mathbf{l}|_E \neq \mathbf{i}$.

* Sinon, ce circuit contient et des points de E et des points de E' . Quitte à effectuer une permutation circulaire des indices, on peut supposer que ce circuit s'écrit $x_1, x_2, \dots, x_n, x_1$ avec $x_1 \in E$ et $x_2 \in E'$. Appelons i le plus petit indice > 2 tel que x_i soit dans E . Comme $\mathbf{l}$ est un ordre strict et que le circuit choisi est de longueur minimale, x_i et x_1 sont distincts, ainsi que x_2 et x_{i-1} . On a alors $x_1 \mathbf{l} x_2$ et $x_{i-1} \mathbf{l} x_i$, et, comme le circuit est de longueur minimale $\neg(x_i \mathbf{l} x_1)$ et $\neg(x_{i-1} \mathbf{l} x_2)$, ce qui contredit $\mathbf{T}(\mathbf{l}, E, E')$.

◇

DÉFINITION 1.28. Soient E et E' deux ensembles finis d'intersection $\{x\}$. On dit qu'une relation $\mathbf{l}$ sur $E \cup E'$ satisfait la condition $\mathbf{T}_x(\mathbf{l}, E, E')$ lorsque:

$$\mathbf{T}_x(\mathbf{l}, E, E') : \quad \forall y \in E \quad \forall y' \in E' \quad \left| \begin{array}{l} y \mathbf{l} y' \Rightarrow (y \mathbf{l} x \text{ ou } x \mathbf{l} y') \\ \text{et} \\ y' \mathbf{l} y \Rightarrow (y' \mathbf{l} x \text{ ou } x \mathbf{l} y) \end{array} \right.$$



THÉORÈME 1.29. Soient E et E' deux ensembles finis tels que $E \cap E' = \{x\}$, $\mathbf{i}$ un ordre strict sur E et $\mathbf{i}'$ un ordre strict sur E' , et $\mathbf{l}$ un ordre strict sur $E \cup E'$ contenant $\mathbf{i} \cup \mathbf{i}'$. Les propositions suivantes sont équivalentes:

$$(i) \quad \forall \mathbf{u} \in \mathcal{R}(E) \quad \forall \mathbf{u}' \in \mathcal{R}(E') \quad (\mathbf{u} \perp \mathbf{i} \text{ et } \mathbf{u}' \perp \mathbf{i}') \Rightarrow \mathbf{l} \perp (\mathbf{u} \cup \mathbf{u}')$$

$$(ii) \left\{ \begin{array}{l} |I|_E = i \text{ et } |I|_{E'} = i' \\ \text{et} \\ T(I, E, E' - \{x\}) \text{ et } T(I, E - \{x\}, E') \\ \text{et} \\ T_x(I, E, E') \end{array} \right.$$

DÉMONSTRATION: Montrons que (i) implique (ii).

* Il est clair que (i) entraîne $T(I, E, E' - \{x\})$ et $|I|_E = i$, et $T(I, E - \{x\}, E')$ et $|I|_{E'} = i'$. Il suffit en effet d'appliquer deux fois le théorème précédent: une première fois en remplaçant E' et i' par $E' - \{x\}$ et $i'|_{E' - \{x\}}$, et une seconde fois en remplaçant E et i par $E - \{x\}$ et $i|_{E - \{x\}}$.

* Montrons maintenant que $T_x(I, E, E')$. Soient deux points $y \in E$ et $y' \in E'$, que l'on peut supposer distincts de x , tels que $y|y'$ et soient $u = \{(x, y)\}$ et $u' = \{(y', x)\}$. Comme I n'est pas orthogonale à $u \cup u'$ il s'ensuit que i n'est pas orthogonale à u , ou que i' n'est pas orthogonale à u' . Comme i et i' sont transitives, on a alors $y|x$ ou $x|i'y'$ et donc $y|x$ ou $x|y'$.

Montrons que (ii) implique (i) en montrant qu'il est impossible que $\neg(i)$ et (ii) soient simultanément satisfaites. Soit donc $u \in \mathcal{R}(E)$ et $u' \in \mathcal{R}(E')$ deux relations respectivement orthogonales à i et i' , telles que I ne soit pas orthogonale à $u \cup u'$. La relation $(I \cup u \cup u' - \mathcal{I}\mathcal{O})$ contient donc des circuits élémentaires qui ne sont ni des circuits de $I - (u \cup u')$, ni des circuits de $(u \cup u') - I$. Prenons en un c qui soit de longueur minimale. En vertu de la démonstration du précédent théorème, on peut se restreindre au cas où x figure dans ce circuit, et les symétries du problème montre qu'on peut alors supposer que x figure dans $c = x_0, x, x_2, \dots, x_{n-1}, x_0$ de l'une des deux manières suivantes: $x_0 u x$ et $x|x_2$ ou $x_0 u x$ et $x u' x_2$.

* Si $x_0 u x$ et $x|x_2$, alors il existe deux points consécutifs x_i et $x_{i+1[n]}$ de c satisfaisant: $x_i \in E' - \{x\}$, $x_{i+1[n]} \in E - \{x\}$, et $x_i|x_{i+1[n]}$.

▷ si $x_2 = x_i$, comme I est un ordre, on a $x|x_{i+1[n]}$, ce qui contredit le fait que c soit de longueur minimale.

▷ Sinon $T(I, E, E' - \{x\})$ montre que $x|x_0$ ou $x_i|x_2$, et chacune de ces deux propositions contredit le fait que c soit de longueur minimale.

* Si $x_0 u x$ et $x u' x_2$, comme $E - \{x\}$ et $E' - \{x\}$ sont disjoints, et que x n'apparaît qu'une seule fois dans c , il existe un indice i tel que $x_i|x_{i+1[n]}$ avec $x_i \in E'$ et $x_{i+1[n]} \in E$. La condition $T_x(I, E, E')$ montre alors que $x_i|x$ ou $x|x_{i+1}$, et chacune de ces deux propositions contredit le fait que c soit de longueur minimale.

◇

§G. Circuits dans un Graphe et Orthogonalité

On considère ici des graphes simples sans boucles contenant des arcs et des arêtes. De tels graphes correspondent bien évidemment à des relations binaires anti-réflexives. Néanmoins il est une

différence ennuyeuse entre les deux présentations: dire que la relation associée à $\mathcal{G}$ est sans-circuit est strictement plus fort que de dire que $\mathcal{G}$ est sans-circuit. En effet, un usage², qui convient parfaitement à notre étude des réseaux ordonnés, est de ne considérer, dans un graphe, que les circuits non-plats: un circuit non-plat d'un graphe est un circuit de la relation qui ne contient pas (x, y) et (y, x) . Plus précisément:

DÉFINITION 1.30. On appelle graphe (simple) $\mathcal{G}$ de sommets S , et d'arcs α le couple (S, α) où α est une relation binaire antiréflexive sur S . Les chaînes de α et les suites d'arcs consécutifs de $\mathcal{G}$ — appelées aussi chemins, pour des raisons évidentes — se correspondent bijectivement.

Un chemin $c : a_0 = (x_0, x_1), a_1 = (x_1, x_2), \dots, a_{n-1} = (x_{n-1}, x_n)$ d'un graphe $\mathcal{G} = (S, \alpha)$ tel que $x_n = x_0$ est appelé un circuit de $\mathcal{G}$ si et seulement si

$$a \in c \Rightarrow \bar{a} \notin c$$

ou

$$\bar{a} = (y, x) \iff a = (x, y)$$

DÉFINITION 1.31. Soit $(\mathcal{G}, \alpha)$ un graphe simple; on appelle connexité de $\mathcal{G}$ la relation $\bar{\alpha}$; autrement dit, $x \bar{\alpha} y$ si et seulement s'il existe un chemin de x à y dans $\mathcal{G}$. On notera que, le graphe n'étant pas symétrique, la connexité n'est pas forcément symétrique, et n'est donc pas, en général, une relation d'équivalence. La plus petite relation d'équivalence qui la contienne s'appelle la connexité faible, et la plus grande relation d'équivalence qu'elle contienne s'appelle la connexité forte.

Etant donné deux graphes sans circuit $\mathcal{G}_1 = (S_1, \alpha_1)$ et $\mathcal{G}_2 = (S_2, \alpha_2)$, une question vient naturellement à l'esprit: leur réunion $\mathcal{G}_1 \cup \mathcal{G}_2 = (S_1 \cup S_2, \alpha_1 \cup \alpha_2)$ est-elle aussi sans circuit? Le théorème suivant, et surtout son corollaire montre qu'il est nul besoin de connaître explicitement α_1 et α_2 pour répondre à cette question. Ces résultats nous seront fort utiles pour établir la modularité des réseaux ordonnés.

THÉORÈME 1.32. Soient $\mathcal{G}_1 = (S_1, \alpha_1)$ et $\mathcal{G}_2 = (S_2, \alpha_2)$ deux graphes simples sans circuit. On a le résultat suivant:

$$\alpha_1 \perp \alpha_2 \Rightarrow \mathcal{G}_1 \cup \mathcal{G}_2 = (S_1 \cup S_2, \alpha_1 \cup \alpha_2) \text{ est sans circuit}$$

De plus, si $\mathcal{G}_1$ et $\mathcal{G}_2$ satisfont la condition:

$$C: \quad \forall u \ (u \in \alpha_1 \Rightarrow \bar{u} \notin \alpha_2) \wedge (u \in \alpha_2 \Rightarrow \bar{u} \notin \alpha_1)$$

alors

$$\alpha_1 \perp \alpha_2 \iff \mathcal{G}_1 \cup \mathcal{G}_2 = (S_1 \cup S_2, \alpha_1 \cup \alpha_2) \text{ est sans circuit}$$

DÉMONSTRATION: Si $\mathcal{G}_1 \cup \mathcal{G}_2$ contient un circuit, il en contient un c qui est élémentaire. Comme $\mathcal{G}_1$ et $\mathcal{G}_2$ sont sans circuit, c contient un arc de $\alpha_1 - \alpha_2$ et un arc de $\alpha_2 - \alpha_1$. En vertu de la caractérisation de l'orthogonalité établie par la proposition 1.24., les relations α_1 et α_2 ne sont pas orthogonales.

²qui correspond par exemple à l'étude homologique d'un graphe considéré comme un complexe simplicial de dimension un

Supposons maintenant que $\mathcal{G}_1$ et $\mathcal{G}_2$ satisfasse la condition $\mathbb{C}$ et que $\mathbf{a}_1$ et $\mathbf{a}_2$ ne soient pas orthogonales. Il existe alors un circuit élémentaire c de $\mathbf{a}_1 \cup \mathbf{a}_2 - \mathcal{I}\mathcal{D}$ qui contient un arc de $\mathbf{a}_1$ et un arc de $\mathbf{a}_2$. Si ce circuit contenait deux arcs u et $\tilde{u}$, alors il serait de la forme x_1, x_2, x_1 avec $u = (x_1, x_2)$ et $\tilde{u} = (x_2, x_1)$, et la condition $\mathbb{C}$ ne serait pas satisfaite. Il s'ensuit que c est un circuit de $\mathcal{G}_1 \cup \mathcal{G}_2$. $\diamond$

COROLLAIRE 1.33. Soient $\mathcal{G}_1 = (S_1, \mathbf{a}_1)$ et $\mathcal{G}_2 = (S_2, \mathbf{a}_2)$ deux graphes simples sans circuit. Alors les deux propositions suivantes sont équivalentes:

- (i) $\left| \begin{array}{l} \mathbb{C}' : \quad \forall u \quad \left(u \in \mathbf{a}_1|_{S_1 \cap S_2} \Rightarrow \tilde{u} \notin \mathbf{a}_2|_{S_1 \cap S_2} \right) \wedge \left(u \in \mathbf{a}_2|_{S_1 \cap S_2} \Rightarrow \tilde{u} \notin \mathbf{a}_1|_{S_1 \cap S_2} \right) \\ \text{et} \\ \mathcal{G}_1 \cup \mathcal{G}_2 \text{ sans circuit} \end{array} \right|$
- (ii) $\overline{\mathbf{a}}_1|_{S_1 \cap S_2} \perp \overline{\mathbf{a}}_2|_{S_1 \cap S_2}$

DÉMONSTRATION: En effet la condition $\mathbb{C}'$ est équivalente à la condition $\mathbb{C}$ et d'autre part, $\overline{\mathbf{a}}_1|_{S_1 \cap S_2} \perp \overline{\mathbf{a}}_2|_{S_1 \cap S_2}$ équivaut à $\mathbf{a}_1 \perp \mathbf{a}_2$. En effet si $\mathbf{a}_1 \perp \mathbf{a}_2$ alors $\overline{\mathbf{a}}_1 \perp \overline{\mathbf{a}}_2$ et par stabilité de l'orthogonalité par sous-relation, on en déduit que: $\overline{\mathbf{a}}_1|_{S_1 \cap S_2} \perp \overline{\mathbf{a}}_2|_{S_1 \cap S_2}$. Supposons maintenant que $\mathbf{a}_1$ et $\mathbf{a}_2$ ne soient pas orthogonales. Alors la proposition 1.24. montre qu'il existe une suite $x_0, x_1, \dots, x_{n-1}$ d'éléments de $S_1 \cup S_2$ tels que:

$$x_0(\overline{\mathbf{a}}_1 - \mathcal{I}\mathcal{D})x_1, x_1(\overline{\mathbf{a}}_2 - \mathcal{I}\mathcal{D})x_2, \dots, x_{n-1}(\overline{\mathbf{a}}_1 - \mathcal{I}\mathcal{D})x_n, x_n(\overline{\mathbf{a}}_2 - \mathcal{I}\mathcal{D})x_0$$

Il est clair que les points x_i sont des points de $S_1 \cap S_2$, et que le circuit ci-dessus est un circuit de $\overline{\mathbf{a}}_1|_{S_1 \cap S_2} \cup \overline{\mathbf{a}}_2|_{S_1 \cap S_2} - \mathcal{I}\mathcal{D}$, ce qui montre que $\overline{\mathbf{a}}_1|_{S_1 \cap S_2}$ et $\overline{\mathbf{a}}_2|_{S_1 \cap S_2}$ ne sont pas orthogonales. $\diamond$

Chapitre 2

Agrégats

§A. Introduction

Nous étudions ici les agrégats qui sont des graphes moins spécifiques (et donc plus maniables) que les réseaux, dont les propriétés combinatoires sont cependant les mêmes — comme nous le verrons au chapitre 4. Les agrégats sont une généralisation particulière des graphes dans laquelle une arête est une paire d'ensembles de sommets (au lieu d'une paire de sommets); on remarquera qu'une "arête" généralisée en ce sens a toujours *deux* sommets, ce qui n'est pas le cas des hyperarêtes d'un hypergraphe.

Plus précisément, un agrégat est un graphe dont les arêtes sont colorées, tel que, pour toute couleur α , le sous-graphe de couleur α est biparti-complet. On s'intéresse alors aux chemins bigarrés i.e. aux chemins qui ne contiennent pas deux arêtes de la même couleur. Si l'on pense les agrégats comme des graphes généralisés, ces chemins sont donc l'analogue des chemins simples (ceux qui n'empruntent pas deux fois la même arête). Comme les réseaux correspondent aux agrégats sans cycle bigarré, qui généralisent les forêts, ceux-ci nous intéresseront plus particulièrement. Cela n'est pas trop surprenant: les preuves habituelles sont effectivement des forêts, et un réseau est, grosso modo, une classe d'équivalence de preuves habituelles.

Une couleur α sera dite scindante si la suppression des arêtes de couleur α déconnecte totalement les deux parties du graphe biparti complet de couleur α . En poursuivant l'analogie, il s'agit donc d'un isthme. Nous démontrons alors le résultat suivant: dans un agrégat sans cycle bigarré, il y a toujours une couleur scindante. Ce résultat est à mettre en parallèle avec le résultat bien connu suivant: dans une forêt, toute arête est un isthme (dans notre cas le "toute" est devenu "une").

En fait, l'énoncé exact du théorème ici démontré est: tout point d'un agrégat est relié par un chemin bigarré à un point d'une couleur scindante, ou est relié par un chemin bigarré à un cycle bigarré.

Ce résultat permet d'unifier les diverses preuves des théorèmes de séquentialisations des réseaux multiplicatifs, et établit l'existence de certains chemins bigarrés, dont notre étude des réseaux fera grand usage.

Ce résultat entraîne trivialement les corollaires suivants pour les réseaux avec la règle de mélange:

- * existence d'un **par** scindant dans un réseau (cf ¶B.1. du chapitre 4)
- * existence d'un **tenseur** scindant dans un réseau dont toutes les conclusions sont des tenseur (cf ¶B.2. du chapitre 4).

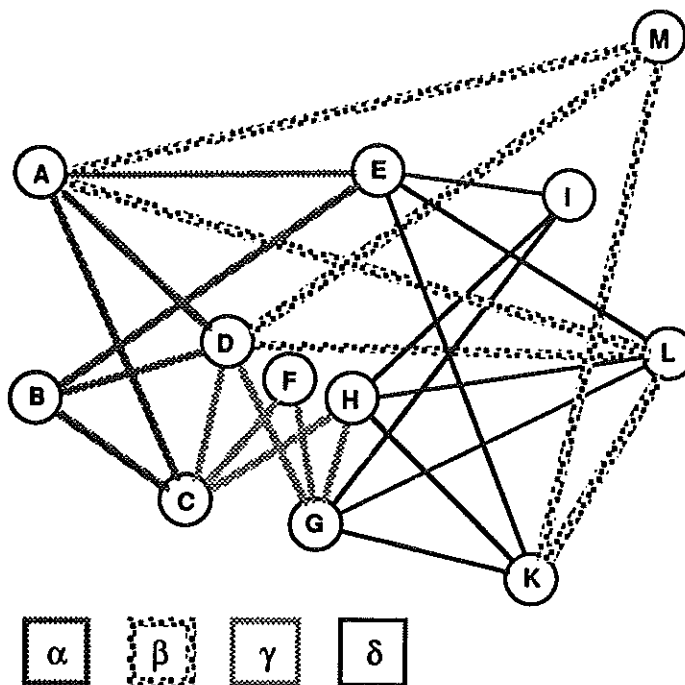
§B. Agrégats

On suivra la terminologie de [Ber79].

On traite ici de graphes dont les arêtes sont colorées. Par commodité, mais en s'éloignant de l'usage, on dira "graphe coloré" pour "graphe dont les arêtes sont colorées".

DÉFINITION 2.1. Un *agrégat* (de graphes bipartis-complets) est un graphe coloré et tel que, pour toute couleur le sous graphe correspondant est un graphe biparti-complet.

Exemple d'agrégat G^0 :



DÉFINITION 2.2. Un ensemble d'arêtes colorées est dit **bigarré** si et seulement s'il ne contient pas deux arêtes d'une même couleur.

DÉFINITION 2.3. Une couleur α d'un agrégat est dite **scindante** ssi les arêtes de couleur α séparent les deux parties du graphe biparti-complet correspondant.

§C. Enoncé des Résultats

Le principal résultat de cette partie est le suivant:

THÉORÈME 2.4. *Soit G un agrégat, et X l'un de ses sommets non isolé ; alors G contient l'une des configurations suivantes:*

$X \xrightarrow{b} \overline{\overline{sc}}$ un chemin bigarré de X à un sommet d'une couleur scindante.

$X \xrightarrow{b} \textcircled{b}$ un chemin bigarré de X à un sommet d'un cycle bigarré.

HYPOTHÈSE 2.5. *On voit que ce résultat vaut si et seulement s'il vaut pour un agrégat connexe; on se limitera donc, dans la suite, aux agrégats connexes.*

Comme le lecteur l'a remarqué nous utilisons différents types de caractères:

- * Les majuscules calligraphiées $\mathcal{G}$ $\mathcal{F}$ $\mathcal{H}$ désignent des agrégats.
- * Les lettres grecques α , β ,... désignent des couleurs.
- * Les majuscules romaines X, Y ,... désignent les sommets.
- * Les minuscules italiques a, b, c ,... désignent des arêtes.
- * Les minuscules sans-serif p, s, c ,... désignent des chemins ou des cycles.

Nous utilisons les notations suivantes:

- * on dira que deux points sont B -connectés s'ils sont reliés par un chemin bigarré
- * $d \in \alpha$ signifie que l'arête d est de couleur α (i.e. est une $\alpha[\mathcal{G}]$ -arête ou, s'il n'y a pas de confusion possible, une α -arête).
- * $\alpha[\mathcal{G}]$ désigne les sommets du sous graphes de couleur α de $\mathcal{G}$. La précision " $[\mathcal{G}]$ " sera éventuellement omise.
- * Parfois, $\alpha[\mathcal{G}]$ désignera les α -arêtes de $\mathcal{G}$, comme dans l'expression $\mathcal{G} - \alpha$; le contexte permettra de faire aisément la différence. La précision " $[\mathcal{G}]$ " sera éventuellement omise, comme dans $\mathcal{G} - \alpha$.
- * On dit que X est un α -sommet de $\mathcal{G}$ ssi $X \in \alpha[\mathcal{G}]$, i.e. ssi X est incident à une $\alpha[\mathcal{G}]$ -arête.
- * On désigne par $\alpha^+[\mathcal{G}]$ et $\alpha^-[\mathcal{G}]$ les deux parties du graphe biparti-complet correspondant. (Ce sont donc des ensembles de sommets.) La précision " $[\mathcal{G}]$ " sera éventuellement ommise.
- * $\mathcal{G} - \alpha$ désigne le sous-graphe obtenu à partir de $\mathcal{G}$ en supprimant les α -arêtes.

§D. Etirement d'un Agrégat

Nous allons ici étudier une opération sur les agrégats dont les propriétés sont nécessaires à la preuve des résultats précédemment énoncés. Cette opération dépend de deux paramètres. Le premier est α^+ , l'une des deux parties du graphe biparti complet correspondant à une couleur α , et le second est une couleur β .

Cette opération consiste à:

- * ajouter une α -arête entre tout sommet de α^+ et tout sommet de β qui n'est pas dans α^- , à moins qu'il y en ait déjà une,
- * supprimer les β -arêtes.

Le résultat de cette opération, appelée étirement et notée $\mathcal{G}_{\alpha+\leftarrow\beta}$ est aussi un agrégat (voir la proposition 2.9. ci-dessous), qui a une couleur de moins β ¹. Si l'on procède à plusieurs étirement consécutifs dont le premier paramètre est le même, le résultat ne dépend pas de l'ordre des étirements. (voir proposition 2.16. ci-dessous), et on peut donc définir cette opération avec un ensemble de couleurs pour second paramètre. Notre preuve (prochain paragraphe) utilise de manière cruciale cette opération avec les paramètres α^+ et $\underline{\beta} = \{\beta_1, \dots, \beta_n\}$ dans le cas suivant: il y a un chemin bigarré entre deux sommets de α^- utilisant les couleurs $\beta_1, \dots, \beta_n$ et aucun des β_i -sommets n'est dans α^+ .

Dans ce cas, la preuve a besoin des résultats suivants:

- * toute couleur scindante de l'agrégat étiré $\mathcal{G}_{\alpha+\leftarrow\underline{\beta}}$ est une couleur scindante de $\mathcal{G}$. (voir les propositions 2.12. et 2.14.).
- * tout chemin bigarré de l'agrégat étiré $\mathcal{G}_{\alpha+\leftarrow\underline{\beta}}$ provient d'un chemin bigarré de $\mathcal{G}$ (voir le lemme 2.17.).

En gardant ceci en tête, nous pouvons maintenant commencer notre étude de l'étirement:

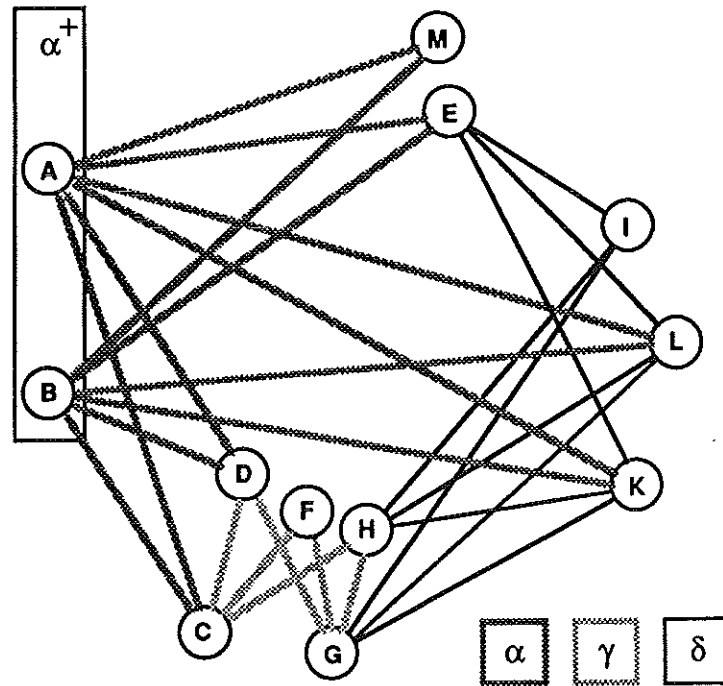
DÉFINITION 2.6. Soient $\mathcal{G}$ un agrégat, α^+ une des deux parties de l'une de ses couleurs et β l'une de ses couleurs. On définit l'étirement $\mathcal{G}_{\alpha+\leftarrow\beta}$ de $\mathcal{G}$ de α^+ vers β comme l'agrégat suivant: (nous le définissons ici comme un graphe coloré, mais la proposition suivante montrera qu'il s'agit bien d'un agrégat):

- * $\mathcal{G}$ et $\mathcal{G}_{\alpha+\leftarrow\beta}$ ont les mêmes sommets.
- * Les couleurs de $\mathcal{G}_{\alpha+\leftarrow\beta}$ sont celles de $\mathcal{G}$ sauf β ²
- * Pour toute couleur $\gamma \neq \alpha, \beta$, $\mathcal{G}_{\alpha+\leftarrow\beta}$ a les mêmes γ -arêtes que $\mathcal{G}$.
- * Il y a une α -arête entre deux sommets de $\mathcal{G}_{\alpha+\leftarrow\beta}$ si et seulement si l'une des deux clauses suivante est satisfaite:
 - ▷ il y en avait déjà une dans $\mathcal{G}$
 - ▷ dans $\mathcal{G}$ l'un des deux sommets était dans $\beta[\mathcal{G}]$ mais pas dans $\alpha^+[\mathcal{G}]$, et l'autre était dans $\alpha^+[\mathcal{G}]$

¹À moins que $\alpha = \beta$, mais nous n'utiliserons pas ce cas trivial

²À moins que $\beta = \alpha$: dans ce cas, $\mathcal{G}_{\alpha+\leftarrow\alpha} = \mathcal{G}$ convient mieux à notre étude

Voici l'étirement $\mathcal{G}_{\alpha+\leftarrow\beta}^0$ de notre exemple $\mathcal{G}^0$:



PROPOSITION 2.7. Pour toute couleur γ de l'étirement $\mathcal{G}_{\alpha+\leftarrow\beta}$ d'un agrégat $\mathcal{G}$, le sous-graphe $\gamma[\mathcal{G}_{\alpha+\leftarrow\beta}]$ est biparti-complet.

DÉMONSTRATION: Pour toute couleur $\gamma \neq \alpha, \beta$ le sous-graphe $\gamma[\mathcal{G}_{\alpha+\leftarrow\beta}]$ de $\mathcal{G}_{\alpha+\leftarrow\beta}$ est biparti-complet puisque c'est le même que $\gamma[\mathcal{G}]$. On vérifie aisément que le sous-graphe $\alpha[\mathcal{G}_{\alpha+\leftarrow\beta}]$ de $\mathcal{G}_{\alpha+\leftarrow\beta}$ est le graphe biparti-complet suivant: l'une de ses deux parties est $\alpha^+[\mathcal{G}_{\alpha+\leftarrow\beta}] = \alpha^+[\mathcal{G}]$ (la même que dans $\mathcal{G}$) et l'autre est $\alpha^-[\mathcal{G}_{\alpha+\leftarrow\beta}] = \alpha^-[\mathcal{G}] \cup (\beta[\mathcal{G}] \setminus \alpha^+[\mathcal{G}])$. Puisque $\alpha^-[\mathcal{G}] \cap \alpha^+[\mathcal{G}] = \emptyset$ les parenthèses peuvent être omises; on les a écrites pour souligner que $\alpha^-[\mathcal{G}_{\alpha+\leftarrow\beta}] \supset \alpha^-[\mathcal{G}]$. $\diamond$

PROPOSITION 2.8. Deux sommets sont connectés dans $\mathcal{G}$ (1) ssi ils le sont dans $\mathcal{G}_{\alpha+\leftarrow\beta}$ (2).

DÉMONSTRATION:

- * (1) $\Rightarrow$ (2) Il suffit de démontrer que si deux sommets étaient adjacents dans $\mathcal{G}$ alors ils sont connectés dans $\mathcal{G}_{\alpha+\leftarrow\beta}$. S'ils sont γ -adjacents avec $\gamma \neq \beta$, alors ils le sont encore dans $\mathcal{G}_{\alpha+\leftarrow\beta}$. S'ils étaient β -adjacents dans $\mathcal{G}$ prenons un sommet Z dans $\alpha^+[\mathcal{G}]$: Z est α -adjacent à l'un et l'autre dans $\mathcal{G}_{\alpha+\leftarrow\beta}$. Ils sont donc connectés dans $\mathcal{G}_{\alpha+\leftarrow\beta}$.
- * (2) $\Rightarrow$ (1) Soient X et Y deux sommets connectés de $\mathcal{G}_{\alpha+\leftarrow\beta}$; ils sont aussi des sommets de $\mathcal{G}$ qui est connexe par 2.5.. $\diamond$

PROPOSITION 2.9. L'étirement $\mathcal{G}_{\alpha+\leftarrow\beta}$ d'un agrégat connexe $\mathcal{G}$ est aussi un agrégat connexe.

DÉMONSTRATION: Conséquence immédiate des deux propositions précédentes. $\diamond$

PROPOSITION 2.10. Les deux graphes $\mathcal{G} - \alpha - \beta$ et $(\mathcal{G}_{\alpha+\leftarrow\beta}) - \alpha$ sont identiques:

DÉMONSTRATION: Evident d'après la définition d'un étirement. $\diamond$

¶D.1. Etirement et Couleurs Scindantes

PROPOSITION 2.11. Si $\gamma \neq \alpha, \beta$ alors les deux agrégats suivants $\mathcal{F}$ et $\mathcal{H}$ sont identiques:

$$\mathcal{F} = (\mathcal{G} - \gamma)_{\alpha+\leftarrow\beta}$$

$$\mathcal{H} = (\mathcal{G}_{\alpha+\leftarrow\beta}) - \gamma$$

DÉMONSTRATION: Remarquons d'abord que ces deux agrégats ont les mêmes sommets (ceux de $\mathcal{G}$) et les mêmes couleurs (celles de $\mathcal{G}$ exceptées γ et β). Il suffit donc de montrer que pour toute couleur $\delta \neq \beta, \gamma$ de $\mathcal{G}$ deux sommets sont δ -adjacents dans $\mathcal{F}$ si et seulement s'ils le sont dans $\mathcal{H}$.

Si $\delta \neq \alpha$, c'est évident: la proposition 2.10. montre que $\mathcal{F} - \alpha = ((\mathcal{G} - \gamma)_{\alpha+\leftarrow\beta}) - \alpha = \mathcal{G} - \gamma - \alpha$ et $\mathcal{H} - \alpha = \mathcal{G}_{\alpha+\leftarrow\beta} - \gamma - \alpha = \mathcal{G}_{\alpha+\leftarrow\beta} - \alpha - \gamma = \mathcal{G} - \alpha - \gamma$.

Comparons maintenant $\alpha[\mathcal{F}]$ et $\alpha[\mathcal{H}]$.

On a $\alpha^+[\mathcal{F}] = \alpha^+[\mathcal{G} - \gamma] = \alpha^+[\mathcal{G}]$ et $\alpha^+[\mathcal{H}] = \alpha^+[\mathcal{G}_{\alpha+\leftarrow\beta}] = \alpha^+[\mathcal{G}]$; ces deux parties sont donc les mêmes.

On sait que $\alpha^-[\mathcal{F}] = \alpha^-[\mathcal{G} - \gamma] \cup (\beta[\mathcal{G} - \gamma] \setminus \alpha^+[\mathcal{G} - \gamma]) = \alpha^-[\mathcal{G}] \cup (\beta[\mathcal{G}] \setminus \alpha^+[\mathcal{G}])$ tandis que $\alpha^-[\mathcal{H}] = \alpha^-[\mathcal{G}_{\alpha+\leftarrow\beta}] = \alpha^-[\mathcal{G}] \cup (\beta[\mathcal{G}] \setminus \alpha^+[\mathcal{G}])$; ces deux parties sont donc aussi les mêmes, ce qui achève la démonstration. $\diamond$

PROPOSITION 2.12. Si $\gamma \neq \alpha$ est une couleur scindante de $\mathcal{G}_{\alpha+\leftarrow\beta}$ alors γ est une couleur scindante de $\mathcal{G}$.

DÉMONSTRATION: Nous allons plutôt montrer que si γ n'est pas une couleur scindante de $\mathcal{G}$, alors γ n'est pas une couleur scindante de $\mathcal{G}_{\alpha+\leftarrow\beta}$. On remarquera que toute couleur de $\mathcal{G}_{\alpha+\leftarrow\beta}$ est une couleur de $\mathcal{G}$, et l'hypothèse signifie donc qu'il existe deux sommets dans des parties opposées de $\gamma[\mathcal{G}]$ qui sont connectés dans $\mathcal{G} - \gamma$. Puisque l'étirement préserve la connexité (proposition 2.8.), ils sont aussi connectés dans $(\mathcal{G} - \gamma)_{\alpha+\leftarrow\beta}$. La proposition 2.11. affirme alors qu'ils le sont aussi dans $(\mathcal{G}_{\alpha+\leftarrow\beta}) - \gamma$. Ainsi ces deux sommets qui sont

bien évidemment dans des parties opposées de γ dans $\mathcal{G}_{\alpha+\leftarrow\beta}$ (puisque $\gamma[\mathcal{G}] = \gamma[\mathcal{G}_{\alpha+\leftarrow\beta}]$) montrent que γ n'est pas scindante dans $\mathcal{G}_{\alpha+\leftarrow\beta}$. $\diamond$

REMARQUE 2.13. La réciproque est vraie: Si $\gamma \neq \alpha, \beta$ est une couleur scindante de $\mathcal{G}$ alors γ est une couleur scindante de $\mathcal{G}_{\alpha+\leftarrow\beta}$.

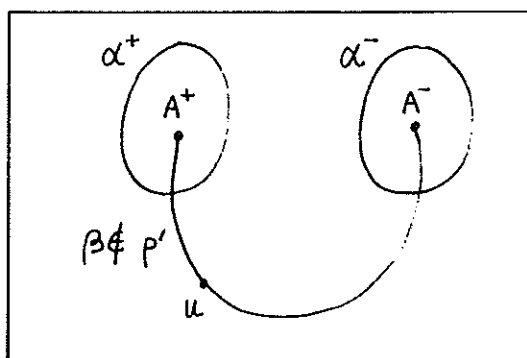
PROPOSITION 2.14. Supposons que $\beta[\mathcal{G}]$ et $\alpha^+[\mathcal{G}]$ soient disjoints.³ Si α est une couleur scindante de $\mathcal{G}_{\alpha+\leftarrow\beta}$ alors α est une couleur scindante de $\mathcal{G}$.

DÉMONSTRATION: Comme précédemment, nous démontrerons que si α n'est pas une couleur scindante de $\mathcal{G}$ alors α n'est pas une couleur scindante de $\mathcal{G}_{\alpha+\leftarrow\beta}$.

Si α n'est pas scindante dans $\mathcal{G}$, il existe $A^+ \in \alpha^+[\mathcal{G}]$ et $A^- \in \alpha^-[\mathcal{G}]$ tels que A^+ et A^- soient connectés dans $\mathcal{G} - \alpha$ par un chemin p .

Si le chemin p ne contient pas de β -arêtes, ces sommets, qui sont aussi dans chacun dans l'une des deux parties de $\alpha[\mathcal{G}_{\alpha+\leftarrow\beta}]$ (puisque $\alpha^+[\mathcal{G}_{\alpha+\leftarrow\beta}] = \alpha^+[\mathcal{G}]$ et $\alpha^-[\mathcal{G}_{\alpha+\leftarrow\beta}] \supset \alpha^-[\mathcal{G}]$), sont aussi connectés dans $\mathcal{G}_{\alpha+\leftarrow\beta} - \alpha = \mathcal{G} - \alpha - \beta$. Ainsi α n'est pas scindante dans $\mathcal{G}_{\alpha+\leftarrow\beta}$.

Si le chemin p contient une β -arête (dessin ci-dessous), il y a un $\beta[\mathcal{G}]$ -sommets sur le chemin p , et soit U le $\beta[\mathcal{G}]$ -sommets le plus près de A^+ sur ce chemin. Alors U est dans $\alpha^-[\mathcal{G}_{\alpha+\leftarrow\beta}]$ (puisque notre hypothèse $\beta[\mathcal{G}] \cap \alpha^+[\mathcal{G}] = \emptyset$ entraîne $\alpha^-[\mathcal{G}_{\alpha+\leftarrow\beta}] = \alpha^-[\mathcal{G}] \cup (\beta[\mathcal{G}] - \alpha^+[\mathcal{G}]) = \alpha^-[\mathcal{G}] \cup \beta[\mathcal{G}] \ni U$). La partie p' de p qui mène à U ne contient ni α -arête ni β -arête, et est donc un chemin du graphe $\mathcal{G} - \alpha - \beta = \mathcal{G}_{\alpha+\leftarrow\beta} - \alpha$. Nous avons donc trouvé deux sommets, l'un dans $\alpha^+[\mathcal{G}_{\alpha+\leftarrow\beta}]$, A^+ , l'autre dans $\alpha^-[\mathcal{G}_{\alpha+\leftarrow\beta}]$, U , qui sont connectés dans $\mathcal{G}_{\alpha+\leftarrow\beta} - \alpha$. Ainsi α n'est pas une couleur scindante de $\mathcal{G}_{\alpha+\leftarrow\beta}$.



$\diamond$

REMARQUE 2.15. La réciproque n'est pas vraie mais on a: si α est une couleur scindante de $\mathcal{G}$, α est une couleur scindante de $\mathcal{G}_{\alpha+\leftarrow\beta}$, à moins qu'un sommet de $\beta[\mathcal{G}]$ soit connecté dans $\mathcal{G} - \alpha$ à un sommet de $\alpha^+[\mathcal{G}]$.

³ Cette	hypothèse	est	bien	nécessaire:	prendre	par	exemple	$\mathcal{G}$	défini	par:
$\alpha^+[\mathcal{G}] =$	$\{X\}$	$\alpha^-[\mathcal{G}] =$		$\{Y\}$						
$\beta^+[\mathcal{G}] =$	$\{X\}$	$\beta^-[\mathcal{G}] =$		$\{Z\}$						
$\gamma^+[\mathcal{G}] =$	$\{Y\}$	$\gamma^-[\mathcal{G}] =$		$\{Z\}$						

¶D.2. Etirement et Chemins Bigarrés

PROPOSITION 2.16. Si l'on effectue deux étirements de même premier paramètre sur un agrégat $\mathcal{G}$ le résultat ne dépend pas de l'ordre dans lequel on les effectue:

$$(\mathcal{G}_{\alpha^+ \leftarrow \gamma})_{\alpha^+ \leftarrow \beta} = (\mathcal{G}_{\alpha^+ \leftarrow \beta})_{\alpha^+ \leftarrow \gamma}$$

DÉMONSTRATION: Ce sont tous les deux des agrégats, et, pour toute couleur $\gamma \neq \alpha, \beta$ le sous-graphe γ de ces deux agrégats est clairement le même, puisque c'est $\gamma[\mathcal{G}]$. Leurs sous-graphes bipartis-complets α ont une partie identique, α^+ tandis que leurs parties α^- sont respectivement $(\alpha^- \cup (\gamma \setminus \alpha^+)) \cup (\beta \setminus \alpha^+)$ et $(\alpha^- \cup (\beta \setminus \alpha^+)) \cup (\gamma \setminus \alpha^+)$, où les couleurs et leurs parties sont prises dans $\mathcal{G}$. Elles sont évidemment égales. $\diamond$

Cette proposition permet de considérer l'étirement d'un agrégat où le second paramètre est un ensemble de couleurs $\underline{\beta} = \{\beta_1, \dots, \beta_n\}$; on notera alors $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$ pour

$$\left(\dots \left((\mathcal{G}_{\alpha^+ \leftarrow \beta_1})_{\alpha^+ \leftarrow \beta_2} \right) \dots \right)_{\alpha^+ \leftarrow \beta_n} = \left(\dots \left((\mathcal{G}_{\alpha^+ \leftarrow \beta_{\varrho(1)}})_{\alpha^+ \leftarrow \beta_{\varrho(2)}} \right) \dots \right)_{\alpha^+ \leftarrow \beta_{\varrho(n)}}$$

pour toute permutation ϱ de $\{1, \dots, n\}$.

Il est alors clair, en écrivant $\mathcal{G} - \underline{\beta}$ pour $\mathcal{G} - \beta_1 - \dots - \beta_n$, que $\mathcal{G} - \underline{\beta} - \alpha = (\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}) - \alpha$.

LEMME 2.17. Supposons que l'on ait des sommets Z_0 et Z_n dans α^- qui soient B-connectés par le chemin bigarré $q = Z_0, e_1 \in \beta_1, Z_1, \dots, Z_{p-1}, e_p \in \beta_p, Z_p, \dots, e_n \in \beta_n, Z_n$ où aucune des couleurs β_i ne soit égale à α .

Notons $\underline{\beta}$ l'ensemble $\{\beta_1, \dots, \beta_n\}$.

alors, si les sommets X et Y sont B-connectés dans $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$, ils le sont B-connectés dans $\mathcal{G}$.⁴

DÉMONSTRATION: Pour éviter les lourdeurs d'écriture on supposera, sauf mention contraire, que les couleurs et leurs parties sont considérées dans $\mathcal{G}$: on écrira donc α au lieu de $\alpha[\mathcal{G}]$.

Appelons p le chemin bigarré entre X et Y dans $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$.

Appelons $q_{0 \dots p-1}$ le sous-chemin de q entre Z_0 et Z_{p-1} et $q_{p \dots n}$ le sous-chemin de q entre Z_p et Z_n :

$$q_{0 \dots p-1} = Z_0, e_1, Z_1, \dots, e_{p-1}, Z_{p-1}$$

$$q_{p \dots n} = Z_p, e_{p+1}, Z_{p+1}, \dots, e_n, Z_n$$

Supposons que les parties des couleurs $\underline{\beta}$ soient nommées de sorte que $Z_{i-1} \in \beta_i^-$ et $Z_i \in \beta_i^+$.

⁴On peut se demander ce qu'il en est de la réciproque. Elle est trivialement fausse: il suffit de considérer l'agrégat suivant à trois sommets X, Y, Y' , et deux couleurs α, β , avec $\alpha^+ = \{X\}$, $\alpha^- = \{Y, Y'\}$ et $\beta^+ = \{Y\}$, $\beta^- = \{Y'\}$. Y et Y' sont α -connectés dans $\mathcal{G}$ mais pas dans $\mathcal{G}_{\alpha^+ \leftarrow \beta} = \mathcal{G} - \beta$.

On remarque que p contient au plus une α -arête puisque c'est un chemin bigarré, et que p ne contient pas de β_i -arête puisque les β_i ne sont pas des couleurs de $\mathcal{G}_{\alpha+\leq \beta}$.

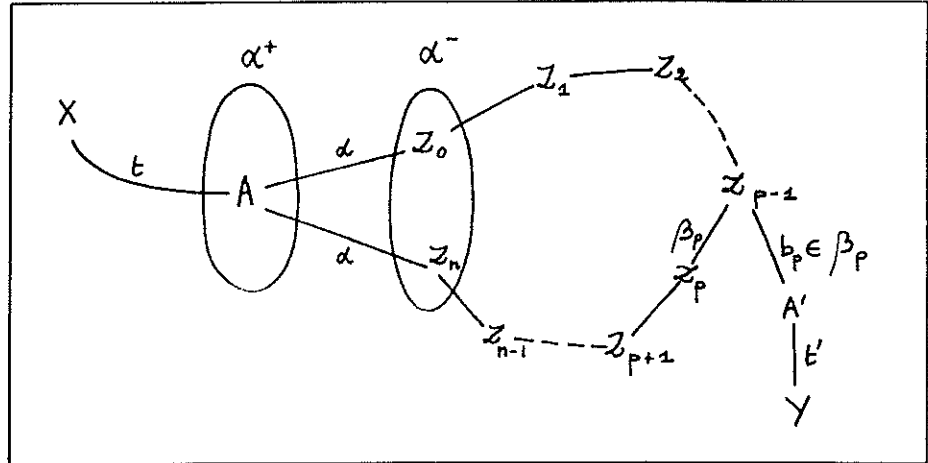
Si p ne contient pas d' α -arête qui ne soit une α -arête de $\mathcal{G}$ il n'y a rien à démontrer: p est un chemin bigarré de $\mathcal{G}$.

Si p contient une α -arête $a = \{A, A'\}$ qui n'est pas une α -arête de $\mathcal{G}$, alors p s'écrit $X, t, A, a \in \alpha, A', t', Y$ où $A \in \alpha^+[\mathcal{G}_{\alpha+\leq \beta}]$ et $A' \in \alpha^-[\mathcal{G}_{\alpha+\leq \beta}]$ (ou l'inverse, mais cela n'a pas d'importance). Donc $A \in \alpha^+$ et il existe une α -arête a_0 de $\mathcal{G}$ entre A et Z_0 , et une α -arête a_n de $\mathcal{G}$ entre A et Z_n . Ainsi $A' \in \beta_p$ pour un certain p , et deux cas symétriques se présentent: $A' \in \beta_p^+$ ou $A' \in \beta_p^-$.

Si $A' \in \beta_p^+[\mathcal{G}]$, il y a dans $\mathcal{G}$ une β_p -arête b_p entre Z_{p-1} et A' . On vérifie aisément que le chemin

$$X, t, A, a_0 \in \alpha, Z_0, s', Z_{p-1}, b_p, A', t', Y$$

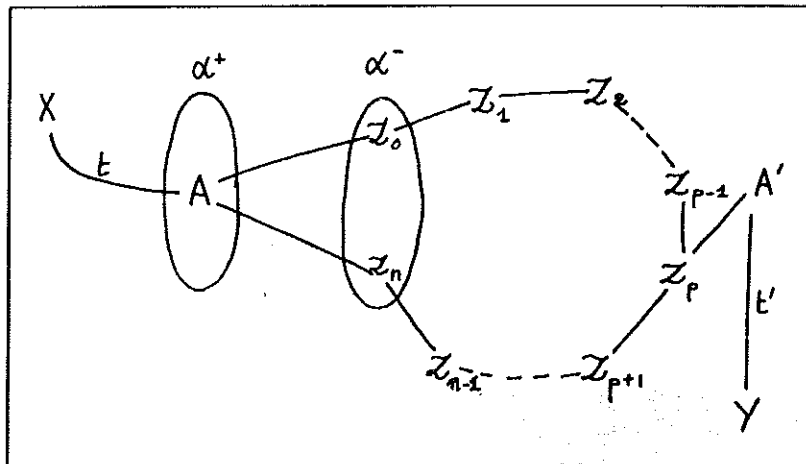
est un chemin bigarré de $\mathcal{G}$. Donc X et Y sont B-connectés dans $\mathcal{G}$. (figure ci-dessous)



Sinon, $A' \in \beta_p^-$, et il y a dans $\mathcal{G}$ une β_p -arête b'_p entre Z_p et A' . On vérifie aisément que le chemin

$$X, t, A, a_n \in \alpha, Z_n, s, Z_p, b'_p \in \beta_p, A', t', Y$$

est un chemin bigarré de $\mathcal{G}$. Donc X et Y sont B-connectés dans $\mathcal{G}$. (figure ci-dessous)



◇

¶D.3. Etirement et Configurations Recherchées.

Dans ce sous-paragraphe, nous supposons que les hypothèses des deux dernières propositions 2.17. 2.14. sont satisfaites:

HYPOTHÈSE 2.18. (pour le reste de ce sous-paragraphe) α^- contient des sommets Z_0 et Z_n qui sont B-connectés par le chemin bigarré $q = Z_0, e_1 \in \beta_1, Z_1, \dots, Z_{p-1}, e_p \in \beta_p, Z_p, \dots, e_n \in \beta_n, Z_n$ et tous β_i sont disjoints de α^+ .

Nous allons démontrer que, si $X \xrightarrow{b} \overset{ec}{\text{---}} \text{---}$ ou $X \xrightarrow{b} \textcircled{b}$ est vrai dans $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$, alors il en est de même dans $\mathcal{G}$.

PROPOSITION 2.19. Une couleur scindante ε de $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$ est couleur scindante de $\mathcal{G}$.

DÉMONSTRATION: On procède par induction sur la construction de $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$:

$$\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}} = \left(\dots \left(\left(\mathcal{G}_{\alpha^+ \leftarrow \beta_1} \right)_{\alpha^+ \leftarrow \beta_2} \right) \dots \right)_{\alpha^+ \leftarrow \beta_n}$$

Il suffit donc de montrer que, pour tout $p \in [1, n-1]$, ε est scindante dans

$$\left(\dots \left(\left(\mathcal{G}_{\alpha^+ \leftarrow \beta_1} \right)_{\alpha^+ \leftarrow \beta_2} \right) \dots \right)_{\alpha^+ \leftarrow \beta_p}$$

entraîne ε est scindante dans

$$\left(\left(\dots \left(\left(\mathcal{G}_{\alpha^+ \leftarrow \beta_1} \right)_{\alpha^+ \leftarrow \beta_2} \right) \dots \right)_{\alpha^+ \leftarrow \beta_p} \right)_{\alpha^+ \leftarrow \beta_{p+1}}$$

Nous l'avons déjà démontré pour $\varepsilon \neq \alpha$ (cf proposition 2.12.). Si $\varepsilon = \alpha$, pour pouvoir appliquer la proposition 2.14. il faut que $\alpha^+ \cap \beta_{p+1}$; cela découle de notre hypothèse 2.18. puisque $\alpha^+[\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}] = \alpha^+[\mathcal{G}]$. $\diamond$

LEMME 2.20. Supposons que X soit B-connecté à un sommet Y d'une couleur scindante $\varepsilon[\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}]$ de $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$ par un chemin bigarré t de $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$, i.e. $X \xrightarrow{b} \overset{ec}{\text{---}} \text{---}$ est vrai dans $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$ avec t , Y et ε . alors X est B-connecté à un sommet Z de la couleur $\varepsilon[\mathcal{G}]$ de $\mathcal{G}$ (qui est aussi scindante dans $\mathcal{G}$) par un chemin bigarré s de $\mathcal{G}$, i.e. $X \xrightarrow{b} \overset{ec}{\text{---}} \text{---}$ est vrai dans $\mathcal{G}$ avec s , Z et ε .

DÉMONSTRATION: Le fait que ε soit aussi scindante dans $\mathcal{G}$ découle de la proposition précédente.

On peut supposer que t ne contient pas de $\varepsilon[\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}]$ -arête: si tel était le cas, on pourrait interrompre t au premier sommet suivi d'une $\varepsilon[\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}]$ -arête.

Si $\varepsilon \neq \alpha$, Y est dans $\varepsilon[\mathcal{G}] = \varepsilon[\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}]$ et le lemme 2.17. fournit un chemin bigarré dans $\mathcal{G}$ entre X et Y .

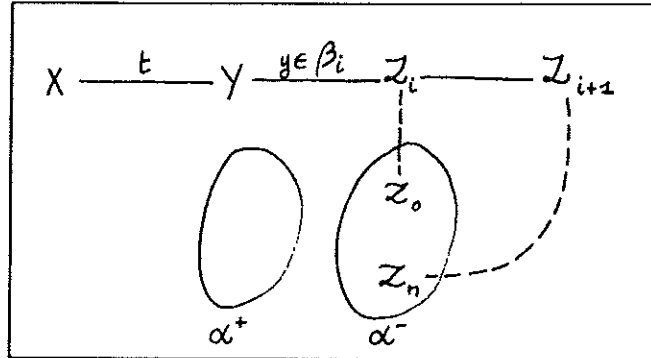
Si $\varepsilon = \alpha$, l'argument ressemble à celui du lemme 2.17.. Puisque t ne contient pas d' α -arête dans $\mathcal{G}_{\alpha^+ \leftarrow \underline{\beta}}$, t est lui-même un chemin bigarré entre X et Y dans $\mathcal{G}$. La seule

difficulté qu'on puisse rencontrer est la suivante: $Y \notin \alpha[\mathcal{G}]$. Cela signifie que $Y \in \beta_i$ pour un $i \in [1, n]$.

Si Y est dans $\beta_i^-[\mathcal{G}]$. Alors il y a une β_i -arête dans $\mathcal{G}$ entre Y et Z_i . On vérifie aisément que le chemin

$$s = X, t, Y, y \in \beta_i, Z_i, e_{i+1}, Z_{i+1}, \dots, Z_n$$

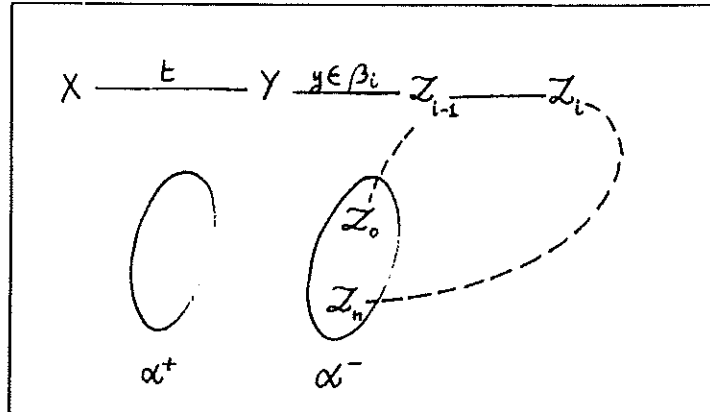
est un chemin bigarré de $\mathcal{G}$ conduisant à Z_n qui appartient à $\alpha[\mathcal{G}]$, i.e. à un sommet de la couleur scindante α de $\mathcal{G}$.



Si Y est dans $\beta_i^+[\mathcal{G}]$. Alors il y a une β_i -arête dans $\mathcal{G}$ entre Y et Z_{i-1} . On vérifie aisément que le chemin

$$s = X, t, Y, y \in \beta_i, Z_{i-1}, e_{i-1}, Z_{i-2}, \dots, Z_0$$

est un chemin bigarré de $\mathcal{G}$ conduisant à Z_0 qui appartient à $\alpha[\mathcal{G}]$, i.e. à un sommet de la couleur scindante α de $\mathcal{G}$.



◇

LEMME 2.21. Supposons que X soit B -connecté à un sommet Y d'un cycle bigarré c de $\mathcal{G}_{\alpha+\leq \beta}$ par un chemin bigarré t de $\mathcal{G}_{\alpha+\leq \beta}$, i.e. $X \xrightarrow{t} Y$ est vrai dans $\mathcal{G}_{\alpha+\leq \beta}$ avec t , Y et c . Alors X est B -connecté à Y du cycle bigarré d de $\mathcal{G}$ par un chemin bigarré s de $\mathcal{G}$, i.e. $X \xrightarrow{s} Y$ est vrai dans $\mathcal{G}$ avec s , Y et d .

DÉMONSTRATION: C'est assez évident, puisque le lemme 2.17. transforme un cycle bigarré c de $\mathcal{G}_{\alpha+\leq \beta}$ en un cycle bigarré d de $\mathcal{G}$ (le cycle transformé contient plus d'arêtes, donc plus de couleurs; il ne peut donc pas être dégénéré).

§E. Démonstration du Résultat Principal

Nous allons maintenant démontrer le théorème 2.4.:

THÉOREME 2.22. (ou 2.4.) Soient G un agrégat, et X l'un de ses sommets non isolé ; alors G contient l'une des configurations suivantes:

$X \xrightarrow{b} \overset{sc}{\Sigma}$ un chemin bigarré de X à un sommet d'une couleur scindante.

$X \overset{b}{\longrightarrow} \textcircled{b}$ un chemin bigarré de X à un sommet d'un cycle bigarré.

E.1. L'Algorithme

L'algorithme procède ainsi:

On va construire de proche en proche un chemin bigarré commençant par X , et lorsque ce ne sera pas possible on recommencera avec X dans un étirement de $\mathcal{G}$, qui a une couleur de moins.

On prend, pour commencer, n'importe quelle arête incidente à X (on traite d'agrégats connexes avec au moins une couleur, c'est donc possible). Désormais, le chemin que nous allons construire commencera toujours par une arête de cette couleur, et ne sera donc jamais vide. On suppose que la dernière arête est

$$Y \in \varepsilon^-, e \in \varepsilon, E \in \varepsilon^+$$

et qu'on est dans l'agrégat $\mathcal{H}$.

On regarde alors s'il y a une arête incidente à un sommet de ε^+ d'une couleur autre que ε .

* S'il n'y en a pas, $\boxed{\text{stop}}$ $X \xrightarrow{b} \overset{sc}{\cancel{X}}$ est vrai dans $\mathcal{H}$. (1)

* S'il y a une arête $d \in \delta$ incidente à E' de ε^+ qui conduit à un point G , on envisage deux cas, suivant que cette couleur δ figure ou non dans le chemin bigarré déjà construit.

▷ Si cette couleur n'a pas été rencontrée, on recommence cette procédure avec le chemin bigarré plus long suivant:

$$\dots(\textit{inchangé})\dots Y, e' \in \varepsilon, E', d \in \delta, G$$

► Si cette couleur δ a déjà été rencontrée: notre chemin bigarré s'écrit $X, t..., U \in \delta^-, v \in \delta, V \in \delta^+, ..s..., Y, e' \in \varepsilon, E'$ et on envisage deux cas, suivant

que $[E \in \delta^- \text{ et } G \in \delta^+] \text{ ou } [E \in \delta^+ \text{ et } G \in \delta^-]$.

- Si $[E \in \delta^- \text{ et } G \in \delta^+]$ alors $\boxed{\text{stop}} X \xrightarrow{b} \textcircled{b}$ est vrai dans $\mathcal{H}$. (2)
- Si $[E \in \delta^+ \text{ et } G \in \delta^-]$ on envisage à nouveau deux cas suivant que l'une des couleurs de s intersecte δ^- ou non.
 - Si l'une des couleurs de s intersecte δ^- , $\boxed{\text{stop}} X \xrightarrow{b} \textcircled{b}$ est vrai dans $\mathcal{H}$. (3)
 - Si aucune couleur de s ne rencontre δ^- on recommence la procédure avec le même sommet X dans l'agrégat $\mathcal{G}_{\delta^- \leftarrow \underline{\beta}}$ (où $\underline{\beta}$ désigne les couleurs de s) et le chemin bigarré

$$X, t, \dots D \in \delta^-[\mathcal{G}_{\delta^- \leftarrow \underline{\beta}}], e \in \delta[\mathcal{G}_{\delta^- \leftarrow \underline{\beta}}], D \in \delta^+[\mathcal{G}_{\delta^- \leftarrow \underline{\beta}}]$$

On remarque que la couleur de la première arête est préservée au cours de cette procédure: même si on étire l'agrégat suivant cette couleur, le chemin se réduit beaucoup mais cela reste vrai, et sinon on garde t et cela reste vrai. Puisque l'on a commencé cette procédure avec un chemin de longueur au moins un, X n'est pas isolé dans $\mathcal{G}_{\delta^- \leftarrow \underline{\beta}}$ (ce qu'assurait déjà la proposition 2.8.).

¶E.2. Correction

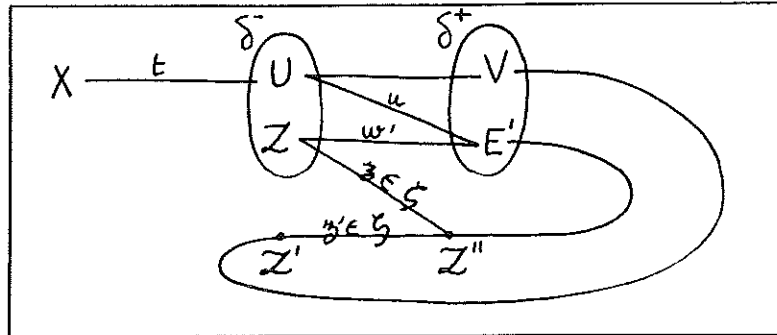
Quand l'agrégat n'a qu'une couleur, $X \xrightarrow{b} \textcircled{b}$ est vrai. A chaque étape de la procédure, soit nous augmentons la longueur du chemin bigarré (et il ne peut y en avoir d'infini), soit nous diminuons le nombre de couleurs de l'agrégat. L'algorithme s'arrête donc en (1), (2) ou (3).

- (1) L'algorithme s'arrête parce que toutes les arêtes incidentes aux sommets de ε^+ sont des ε -arêtes. Dans ce cas il est clair que la couleur ε est scindante dans $\mathcal{H}$. On procède ensuite par induction sur le nombre p d'étirements effectués pour obtenir $\mathcal{H}$ à partir de $\mathcal{G}$, et p applications du lemme 2.20. montrent que $X \xrightarrow{b} \textcircled{b}$ est vrai dans l'agrégat $\mathcal{G}$ de départ. On remarquera que la preuve fournit explicitement la couleur scindante (c'est celle sur laquelle s'arrête l'algorithme), le sommet de la couleur scindante (en utilisant le lemme 2.20.), et le chemin bigarré (en utilisant le lemme 2.17.).
- (2) Le cycle bigarré est $V, \dots s, \dots D, e' \in \varepsilon, E', d' \in \delta, V$ et le chemin bigarré $X, t, \dots, U \in \delta^-, v \in \delta, V$. Ce chemin et ce cycle bigarré peuvent être effectivement transformés en un chemin bigarré de X à V et un cycle bigarré passant par V (en utilisant le lemme 2.21.). Donc $X \xrightarrow{b} \textcircled{b}$ est vrai dans $\mathcal{G}$.
- (3) Si une couleur de s , disons ζ , intersecte δ^- en Z . montrons que $X \xrightarrow{b} \textcircled{b}$ est vrai dans $\mathcal{H}$, et donc dans $\mathcal{G}$ d'après le lemme 2.21.. Dans ce cas, on a le chemin bigarré suivant:

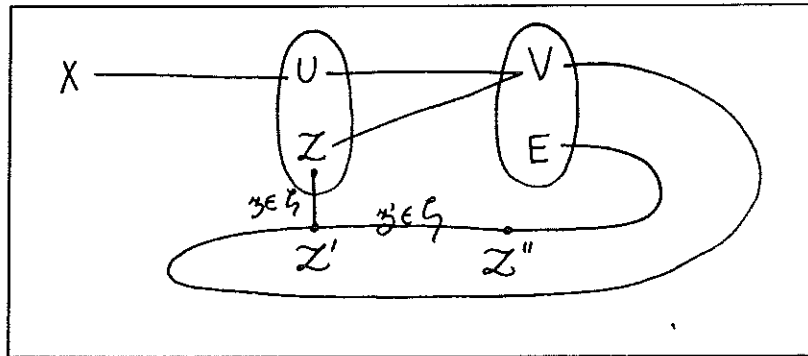
$$X, t, \dots, U \in \delta^-, v \in \delta, V \in \delta^+, \dots Z', z' \in \zeta, Z'' \dots Y, e' \in \varepsilon, E', d \in \delta, G$$

Remarquons qu'on a alors une δ -arête u entre E' et U une δ -arête w entre Z et V , et une δ -arête w' entre Z et E' . Deux cas se présentent:

- * Z et Z' sont dans la même partie de ζ , et il y a donc une ζ -arête z entre Z et Z'' . Le point est E' , le chemin bigarré $X, t, \dots, U \in \delta^-, u \in \delta, E' \in \delta^+$, et le cycle bigarré $E' \in \delta^+, w' \in \delta, Z \in \delta^-, z \in \zeta, Z'' \dots Y, e' \in \varepsilon, E'$



* Z et Z'' sont dans la même partie de ζ , il y a donc une ζ -arête z entre Z et Z' . Le point est V , le chemin bigarré $X, t, \dots, U \in \delta^-, v \in \delta, V \in \delta^+$ et le cycle bigarré $V \in \delta^+, \dots, Z', z \in \zeta, Z, w \in \delta, V$



¶E.3. Complexité de l'Algorithme en l'Absence de Cycle Bigarré

Dans ce cas, qui est celui dont on se sert pour les réseaux, l'algorithme trouve un chemin bigarré de X à un sommet d'une couleur scindante et le fait en un temps raisonnable, puisque la seule longue étape est de vérifier si l'une des couleurs β intersecte δ^+ . Nous pouvons alors donner une estimation du temps nécessaire au calcul de la configuration $X \xrightarrow{b} \text{---} \xrightarrow{\delta^c}$, en appelant C le nombre total de couleurs de l'agrégat de départ.

On choisit au hasard une arête incidente à un sommet de δ^+ , et on teste si cette couleur figure dans le chemin bigarré. Cela prend ℓ étapes où ℓ désigne la longueur du chemin bigarré déjà construit. Au pire, on fait cela jusqu'à ce que ℓ soit maximal, i.e. $\ell = c$. Le nombre de telles étapes est au pire:

$$\sum_{k=1}^{k=c} k = \frac{(c+1)c}{2}$$

Alors l'algorithme s'arrête, ou on recommence dans un agrégat ayant une couleur de moins, et avec un chemin peut-être de nouveau réduit à une arête. Au pire on recommence cela C fois. Ainsi le nombre maximal d'étapes de cet algorithme est moins de:

$$\sum_{c=1}^{c=C} \frac{(c+1)c}{2} = \frac{1}{2} \frac{(2C+1)(C+1)C}{6} + \frac{1}{2} \frac{(C+1)C}{2} \approx \frac{C^3}{6}$$

§F. Corollaires

Pour commencer, on remarque que les couleurs scindantes donne une structure d'arbre à un agrégat connexe, d'où:

DÉFINITION 2.23. La suppression d'une couleur scindante dans un agrégat connexe définit deux classes de sommets: ceux qui sont connectés à un coté, et ceux qui sont connectés à l'autre. On appelle feuille d'un agrégat connexe une telle classe si elle est sans couleur scindante. On appelle feuille d'un agrégat les feuilles de ses composantes connexes. On remarquera qu'une feuille ne rencontre qu'au plus une couleur scindante.

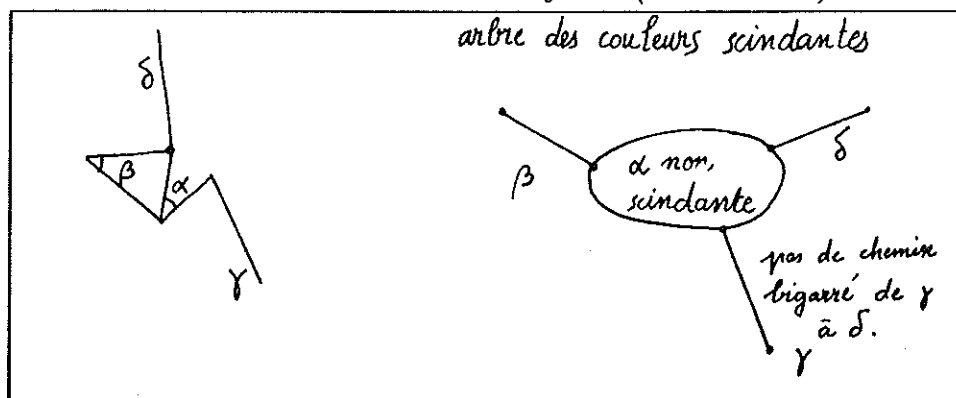
On déduit immédiatement du théorème 2.4. le résultat suivant:

COROLLAIRE 2.24. Si l'agrégat est sans cycle bigarré, alors il existe un chemin bigarré de tout point d'une feuille à la couleur scindante qui la définit.

THÉORÈME 2.25. Soit G un agrégat sans cycle bigarré, et soit α^+ une partie d'une couleur scindante α de G telle qu'il existe d'autres couleurs scindantes du coté α^+ de α . Alors il existe un chemin bigarré de α^+ à l'une de ces couleurs scindantes.

DÉMONSTRATION: On procède par récurrence sur le nombre de couleurs (scindantes ou non-scindantes) situées du coté de α^+ — disons à gauche, pour simplifier — dans G . On suppose qu'il y a au moins une couleur scindante à gauche de α . S'il n'y a qu'une couleur à gauche de α , elle est scindante et le résultat vaut. Sinon on applique l'algorithme en commençant par une arête de α . Si l'on effectue un étirement, on remarque que celui-ci ne détruit que des couleurs non-scindantes à gauche de α , et que les couleurs scindantes situées à gauche de α restent des couleurs scindantes situées à gauche de α . L'hypothèse d'induction nous fournit donc un chemin bigarré de α à une couleur scindante située à gauche de α dans l'étirement et le lemme 2.20. nous permet de construire une configuration similaire dans l'agrégat non-étiré. $\diamond$

REMARQUE 2.26. Il est, par contre, faux qu'il existe un chemin bigarré d'une couleur scindante α à toute couleur scindante située immédiatement à sa gauche (ou à sa droite):



Les deux résultats précédents entraînent immédiatement le suivant:

COROLLAIRE 2.27. Soit G un agrégat sans cycle bigarré, et F_G l'une de ses feuilles. Alors il existe une autre feuille F'_G telle qu'il existe un chemin bigarré entre tout sommet de F_G et tout sommet de F'_G .

§G. Sous-Graphes Bigarrés Maximaux

THÉORÈME 2.28. *Soient G et G' deux sous-graphes bigarrés maximaux d'un agrégat G sans cycle bigarré. Alors G et G' ont le même nombre de composantes connexes.*

DÉMONSTRATION: De tels sous-graphes sont obtenus en choisissant, pour toute couleur α une arête de couleur α , et ce sont des forêts: ils contiennent donc au plus un chemin minimal entre deux sommets. Vu que le nombre de couleurs est fini, G' est un transformé de G par une suite finie de transformations du type suivant:

$$G \longrightarrow G - a + a' \text{ où } a \text{ et } a' \text{ sont deux arêtes de couleur } \alpha$$

Une récurrence triviale montre qu'il suffit de vérifier qu'une transformation de ce type préserve le nombre de composantes connexes. On remarquera qu'une telle transformation ne modifie que les composantes connexes des sommets de $a = XY$ et $a' = X'Y'$.

Montrons d'abord qu'il suffit de traiter le cas où a et a' ont un sommet commun. Supposons que X et X' soient dans la même partie de la couleur α , et donc que Y et Y' soient dans l'autre. Appelons b l' α -arête XY' et b' l' α -arête $X'Y$. On remarquera qu'une telle transformation se décompose en:

$$G \longrightarrow G - a + b \longrightarrow G - b + b' \longrightarrow G - b' + a'$$

qui ne fait intervenir que des transformations où les deux arêtes ont un sommet commun.

Supposons maintenant que a et a' aient un sommet commun, disons $X = X'$. Envisageons deux cas, suivant que le nombre de composantes connexes de G susceptibles d'être modifiées ($cc(X) = cc(Y)$ et $cc(Y')$) est un ou deux.

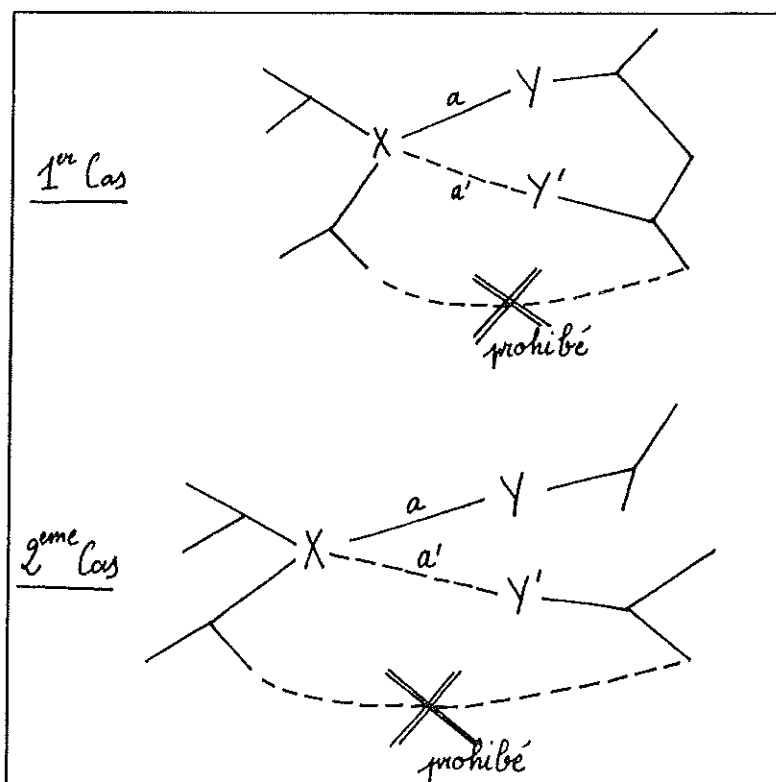
Une seule composante connexe risque d'être modifiée, i.e. $cc(X) = cc(Y) = cc(Y')$ dans G . Remarquons que tout chemin de G reliant X et Y' utilise a : sinon, ce chemin et l'arête a' constitueraient un cycle bigarré de l'agrégat. C'est donc que Y et Y' sont reliés par un chemin de G n'utilisant pas a . L'unique composante connexe susceptible de modification ne s'est pas divisée au cours de la transformation puisqu'on a:

$$cc(X) = cc(Y') = cc(Y) \text{ dans } G'$$

Deux composantes connexes risquent d'être modifiées, i.e. $cc(X) = cc(Y) \neq cc(Y')$ dans G . Remarquons de suite qu'elles ne peuvent se séparer en au moins trois composantes connexes, par suppression d'une arête (et a fortiori par le remplacement d'une arête). Remarquons que tout chemin de G reliant X et Y utilise a : sinon, ce chemin et l'arête a constitueraient un cycle bigarré de l'agrégat. Il n'y a donc pas de chemin entre Y' et Y dans G' : si tel était le cas, ce chemin utiliserait a' , puisqu'il n'existait pas dans G ; il y aurait alors un chemin entre X et Y dans G' , et donc un chemin de G n'utilisant pas a entre X et Y , ce que l'on vient de réfuter. Les deux composantes connexes ont été modifiées mais sont toujours deux puisqu'on a:

$$cc(X) = cc(Y') \neq cc(Y) \text{ dans } G'$$

◇



Partie B

La règle de mélange

Chapitre 3

Séquents, Réseaux et Règle de Mélange

Dans cette partie, nous étudions brièvement le calcul linéaire multiplicatif (dont les seuls connecteurs sont le **tenseur** et le **par** , avec une règle supplémentaire dite de mélange:

$$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \\ \vdots \end{array} \vdash A_1, A_2, \dots, A_n \quad \begin{array}{c} \vdots \pi_2 \\ \vdots \\ \vdots \end{array} \vdash B_1, B_2, \dots, B_p}{\vdash A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p} \text{MÉLANGE}$$

Ce calcul n'est pas le centre de notre travail (qui est le calcul ordonné) mais nous présentons tout de même nos travaux le concernant car:

- ★ Ce calcul est sous-jacent au calcul ordonné, i.e. le cas très élémentaire où l'ordre des conclusions est toujours vide au cours de la preuve.
- ★ C'est l'occasion d'une répétition générale, où présenter les séquents, réseaux et espaces cohérents, et établir les résultats qui les lient dans ce cas plus simple.
- ★ Les agrégats et autres notions combinatoires de la première partie sont utilisées ici pour démontrer certaines propriétés des réseaux qui seront bien utiles lors de l'étude des réseaux ordonnés.
- ★ Ce calcul permet de plus de traiter les éléments neutres (et donc l'affaiblissement) comme on le verra au chapitre 5.

§A. Le Calcul des Séquents

Donnons les règles de ce calcul:

Règle structurelle	
Règle de mélange	$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \end{array} \quad \frac{\vdots \pi_1 \quad \vdots \pi_2}{\vdash A_1, A_2, \dots, A_n \quad \vdash B_1, B_2, \dots, B_p} \text{MÉLANGE}}{\vdash A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p}$
Règles d'identité	
Axiome	$\frac{}{\vdash A, A^\perp} \text{AXIOME}$
Règle de coupure	$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \end{array} \quad \frac{\vdots \pi_1 \quad \vdots \pi_2}{\vdash A_1, A_2, \dots, A_n, K \quad \vdash K^\perp, B_1, B_2, \dots, B_p} \text{COUPURE}}{\vdash A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p}$
Règles logiques	
Règle du tenseur	$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \end{array} \quad \frac{\vdots \pi_1 \quad \vdots \pi_2}{\vdash A_1, A_2, \dots, A_n, A \quad \vdash B, B_1, B_2, \dots, B_p} \text{TENSEUR}}{\vdash A_1, A_2, \dots, A_n, A \otimes B, B_1, B_2, \dots, B_p}$
Règle du par	$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \end{array} \quad \frac{\vdots \pi_1 \quad \vdots \pi_2}{\vdash A, B, A_1, A_2, \dots, A_n} \text{PAR}}{\vdash A \wp B, A_1, A_2, \dots, A_n}$

THÉORÈME 3.1. (élimination des coupures) *Toute preuve de ce calcul peut être transformée en une preuve n'utilisant pas la règle COUPURE ; cette réduction est localement confluente et termine toujours.*

DÉMONSTRATION: C'est une conséquence directe du même résultat pour le calcul plus riche du chapitre 5 (cf 5.4.) $\diamond$

THÉORÈME 3.2. *Toute preuve de ce calcul peut être transformée en une preuve n'utilisant que des axiomes atomiques.*

DÉMONSTRATION: C'est une conséquence directe du même résultat pour le calcul plus riche du chapitre 5 (cf §D.) $\diamond$

§B. Les Réseaux Correspondants

On va les définir comme des graphes dont les sommets sont étiquetés par des formules et dont les arêtes sont colorées.

NOTATION 3.3. [couleurs] *On se donne un ensemble dénombrable de couleurs. Parmi ces couleurs, on en distingue une, le noir, notée $\mathcal{N}$ tandis que les autres seront dites propres, et seront désignées par des lettres grecques $\alpha, \beta, \dots$*

DÉFINITION 3.4. *On appelle (abusivement) arbre des sous-formules d'une formule F , un arbre binaire dont les sommets sont étiquetés par des sous-formules de F et les branchements par des*

connecteurs de sorte que, si G et H sont les successeurs de L dans un branchement étiqueté par le connecteur $\heartsuit$, alors $L = G\heartsuit H$. On remarquera qu'on s'autorise ainsi à considérer des arbres des sous-formules "élagués" dont les feuilles ne sont pas toujours des atomes; par contre, si une sous-formule immédiate H , d'une formule $H\heartsuit G$ figure dans l'arbre, l'autre sous-formule immédiate G y figure aussi.

DÉFINITION 3.5. Une famille d'arbres des sous-formules est dite bien colorée si et seulement si:

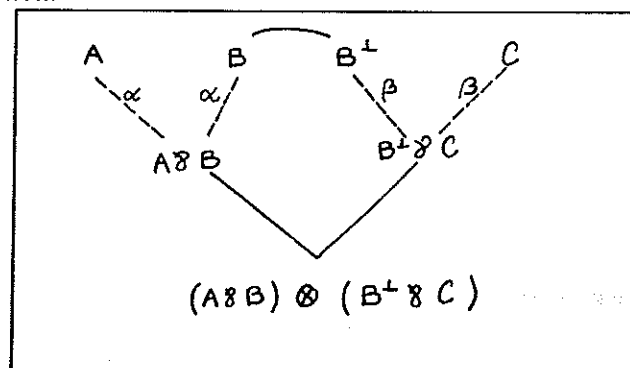
- * les deux arêtes correspondant à un connecteur tenseur sont noires;
- * les deux arêtes correspondant à un connecteur par sont d'une même couleur propre;
- * deux arêtes appartenant à des connecteurs par différents, — qu'ils soient ou non dans le même arbre — sont de couleurs (propres) différentes.

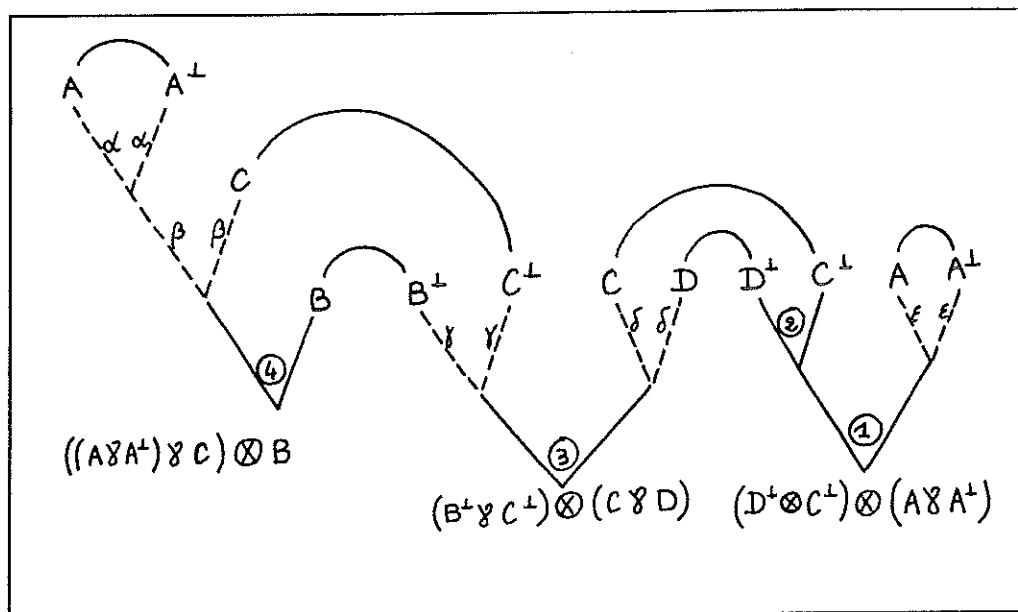
DÉFINITION 3.6. Un pré-réseau est un graphe simple coloré dont les sommets sont étiquetés par des formules et le signe " $\bullet$ ". Plus précisément, un pré-réseau Π est défini par un triplet $(\mathcal{F}, Ax, Cut)$ où:

- * $\mathcal{F}$ est une famille bien colorée d'arbres des sous-formules des formules F_i , dont les feuilles sont les A_i .
- * Ax est un ensemble d'arêtes noires non-adjacentes de la forme $A_i - A_j$ où $A_i = A_j^\perp$
- * Cut est un ensemble triplets $(\bullet_k, c_k, c_k^\perp)$ où:
 - ▷ les $\bullet_k$ sont des sommets du pré-réseau autres que ceux de $\mathcal{F}$, étiquetés $\bullet$, de degré 2, appelés coupures du pré-réseau
 - ▷ pour tout lien k , c_k et $c_k^\perp$ sont des arêtes noires de la forme $\bullet_k - F_i$ et $\bullet_k - F_j$ avec $F_i = F_j^\perp$, de sorte que, si $k \neq k'$, les deux arêtes c_k et $c_{k'}$ ne sont pas adjacentes et les deux arêtes c_k et $c_{k'}$ ne sont pas adjacentes. (On peut aussi dire qu'une racine F_i est incidente à au plus une arête c_k ou $c_k^\perp$)

On appelle *feuille* du pré-réseau toute feuille de la forêt $\mathcal{F}$. On appelle *coupures* du pré-réseau les sommets de Cut (i.e. les $\bullet_k$). On appelle *conclusion* du pré-réseau toute racine F_i de $\mathcal{F}$ non-adjacente à un sommet de Cut (non-incidente à une arête de Cut). On appelle *hypothèse* d'un pré-réseau une feuille non-incidente à une arête de Ax .

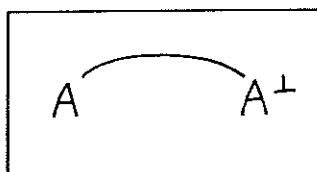
Bien que les pré-réseaux et les agrégats partagent les mêmes propriétés combinatoires, on voit que la définition des premiers est nettement moins maniable. Donnons deux exemples de pré-réseaux afin d'éclaircir cette notion:



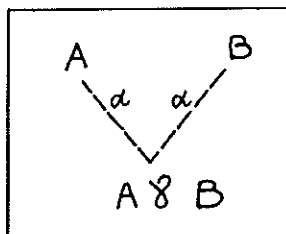


DÉFINITION 3.7. Un lien d'un préréseau $\Pi = (\mathcal{F}, Ax, Cut)$ est l'un de ses (petits) sous-graphes pleins suivants:

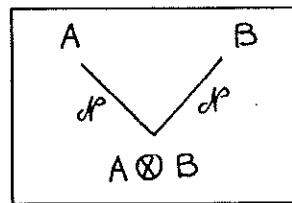
- * Un lien hypothèse est le graphe réduit à une hypothèse du préréseau.
- * Un lien axiome est une arête de Ax . Les deux extrémités de l'arête, A et $A^\perp$ sont dites conclusions de ce lien.



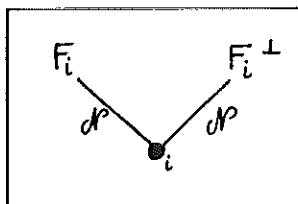
- * Un lien par est constitué des deux arêtes adjacentes d'une même couleur propre liant une sous-formule $A \wp B$ de F_i à ses deux sous-formules immédiates A et B . $A \wp B$ est appelée la conclusion du lien, tandis que ses sous-formules immédiates A et B sont appelées prémisses de ce lien.



- * Un lien tenseur est constitué des deux arêtes noires liant une sous-formule $A \otimes B$ de F_i à ses deux sous-formules immédiates A et B . $A \otimes B$ est appelée la conclusion du lien, tandis que ses sous-formules immédiates A et B sont appelées prémisses de ce lien.



* Un lien coupure est constitué des deux arêtes noires incidentes à un sommet de Cut. Le sommet $\bullet_k$ est dit conclusion de ce lien, tandis que les formule F_k et $F_k^\perp$ sont dites prémisses de ce lien.



DÉFINITION 3.8. Dans un graphe coloré on dit qu'un ensemble d'arêtes est bigarré s'il ne contient pas deux arêtes d'une même couleur propre. Un chemin bigarré est donc un chemin qui n'emprunte jamais deux arêtes d'une même couleur propre — cf Chapitre 2

DÉFINITION 3.9. On dit qu'un préréseau est un réseau s'il ne contient pas de cycle bigarré.

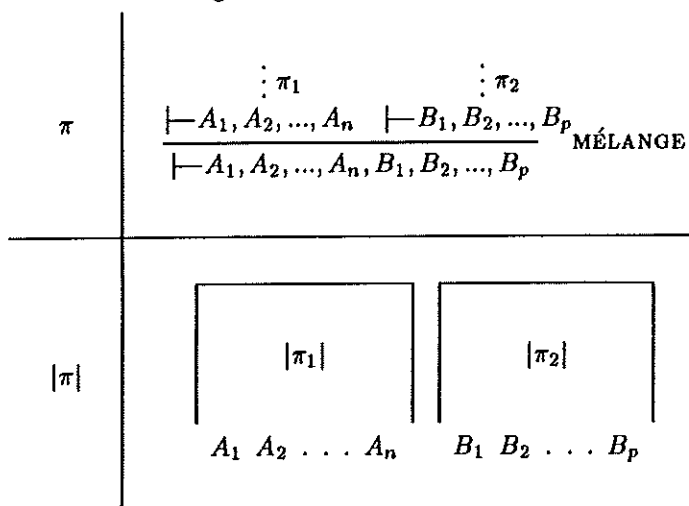
Ainsi le premier exemple n'est pas un réseau, et le second en est un.

§C. Des Preuves Séquentielles aux Réseaux

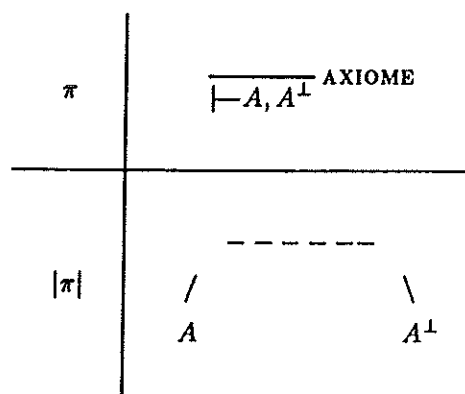
On va associer, à toute preuve π du calcul des séquents précédemment défini, un réseau $|\pi|$ sans hypothèse et dont les conclusions sont celles de π . Pour cela on commence par lui associer un préréseau, et on vérifiera que c'est un réseau sans hypothèse. Cette question étant triviale et bien connue, nous irons un peu vite.

La traduction d'une preuve séquentielle π en un réseau $|\pi|$ se fait inductivement:

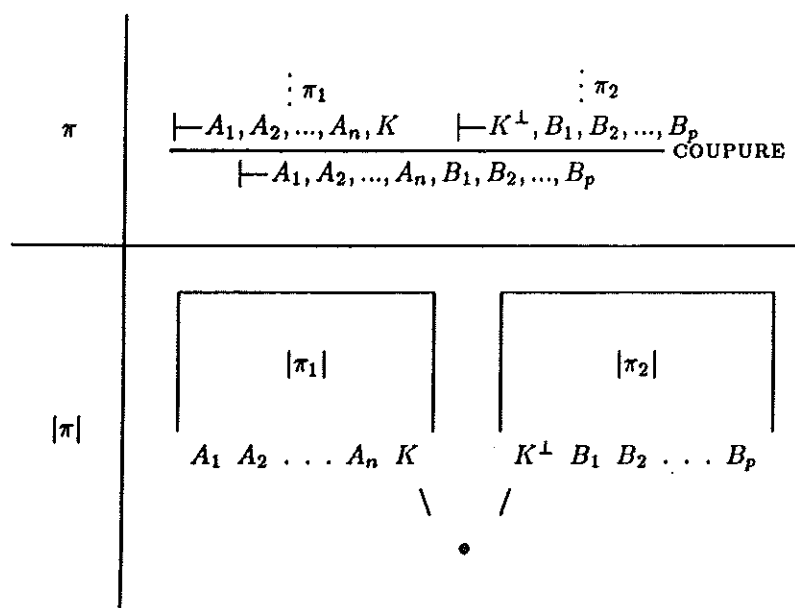
* Si la dernière règle de π est MÉLANGE :



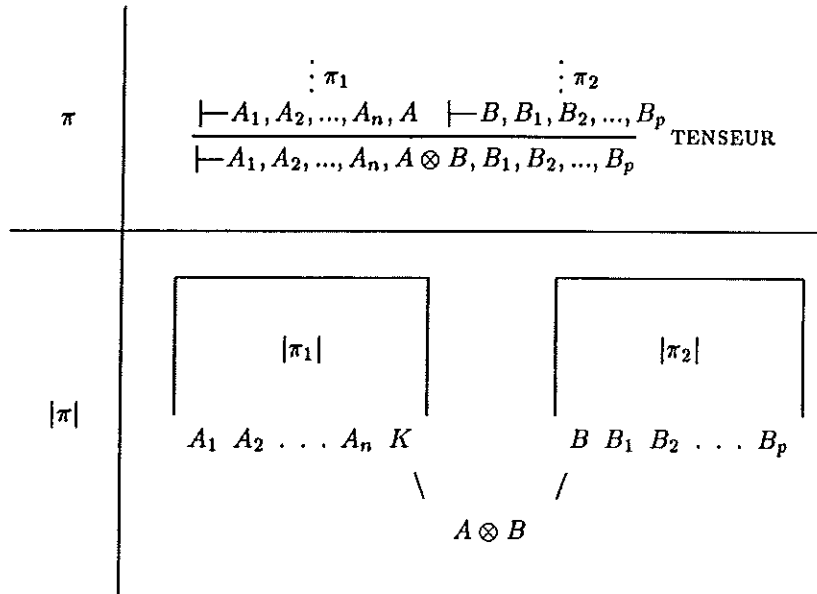
★ Si la dernière règle de π est AXIOME :



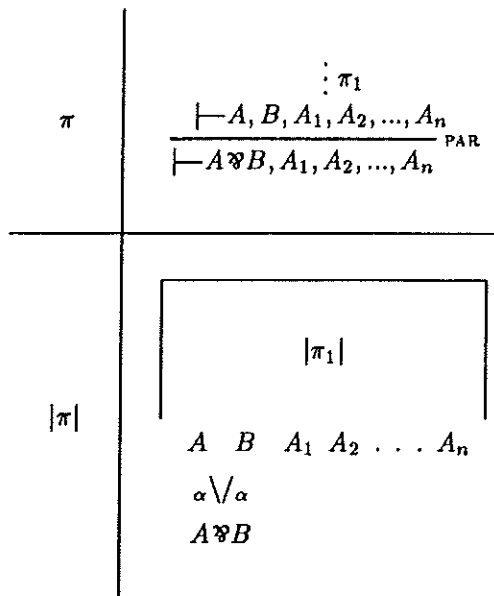
★ Si la dernière règle de π est COUPURE :



★ Si la dernière règle de π est TENSEUR :



★ Si la dernière règle de π est PAR :



PROPOSITION 3.10. *La traduction d'une preuve séquentielle est un réseau sans hypothèse.*

DÉMONSTRATION: Le fait qu'on obtienne ainsi un préréseau sans hypothèse est clair, et le fait qu'il n'y ait pas de cycle bigarré dans la traduction d'une preuve séquentielle se vérifie aisément par induction sur la preuve traduite. $\diamond$

Chapitre 4

Des Réseaux aux Séquents

DÉFINITION 4.1. *On dit qu'un réseau Π est séquentialisable si et seulement s'il existe une preuve π du calcul des séquents telle que le réseau associé à π soit Π . On dit alors que π est une séquentialisation de Π .*

THÉORÈME 4.2. (de séquentialisation) *Tout réseau est séquentialisable.*

Ce chapitre est essentiellement consacré à la démonstration de ce théorème. La preuve repose sur le théorème principal du chapitre précédent 2.4., qu'on applique à deux traductions d'un réseau en agrégat. On unifie ainsi les deux techniques les plus courantes de séquentialisation, qui jusqu'ici étaient sans rapport: la méthode du tenseur scindant de [Gir87a] (qui jusqu'ici ne permettait pas de traiter le calcul avec la règle MÉLANGE) et celle de [Dan90]. L'application des corollaires (chapitre 2 §F.) du théorème 2.4. à ces traductions d'un réseau en agrégat permet de plus d'en déduire l'existence de certains chemins bigarrés dans un réseau, mais cela sera présenté au chapitre 9, lorsque se sera utile pour comprendre la structure des réseaux ordonnés.

On étudie ensuite de manière un peu plus générale le lien entre réseaux et agrégats afin de montrer que réseaux et agrégats sont, d'un point de vue combinatoire, les même objets.

On termine par la notion de tenseur héréditairement scindant, qui nous sera nécessaire pour plonger les réseaux ordonnés dans ces réseaux-ci .

§A. Le théorème de séquentialisation

On appelle hypothèse dans le calcul des séquents un axiome propre:

$$\vdash H$$

où H est une formule. Le réseau correspondant consiste simplement en la formule H . Un réseau correspondant à une preuve utilisant les hypothèses $\vdash H_i$ est un réseau avec hypothèses dont les hypothèses sont précisément les H_i .

Nous allons en fait démontrer le théorème suivant, qui assure évidemment le théorème de séquentialisation:

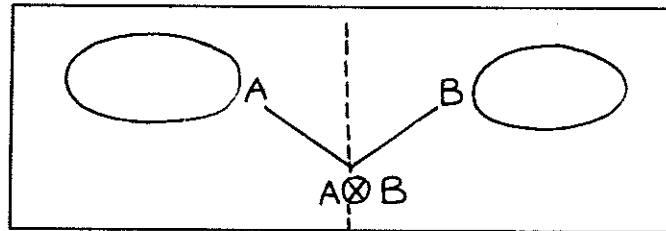
THÉORÈME 4.3. *Pour tout réseau Π de conclusions $C_1, \dots, C_p$ et d'hypothèses $H_1, \dots, H_n$, il existe une preuve séquentielle de $\vdash C_1, \dots, C_n$ sous les hypothèses $\vdash H_1 \dots \vdash H_p$ dont le réseau associé est Π .*

On procède par induction sur le nombre de liens du réseau Π , qu'on appellera la taille de Π . Si ce nombre est 1, le résultat est évident: la preuve correspondant à Π est un axiome ou une hypothèse.

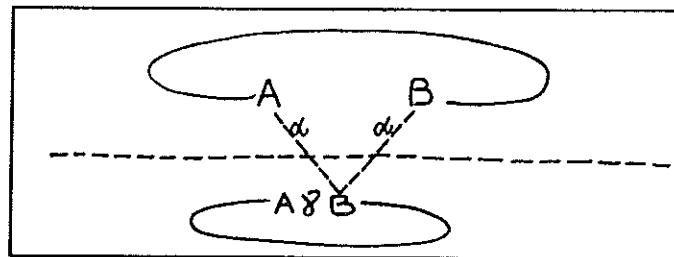
PROPOSITION 4.4. *Si le réseau Π n'est pas connexe, et si tout réseau Π' de taille $l' < l$ est séquentialisable, alors Π est séquentialisable.*

DÉMONSTRATION: Une partition en deux classes des composantes connexes donne deux réseaux Π_1 et Π_2 plus petits (un réseau non vide contient au moins un lien), et la règle MÉLANGE appliquée à une séquentialisation de Π_1 et une séquentialisation de Π_2 donne une séquentialisation de Π . $\diamond$

DÉFINITION 4.5. *Un lien tenseur ou coupure est dit scindant si et seulement s'il est final et si l'une de ses deux arêtes est un isthme (auquel cas l'autre est aussi un isthme).*



DÉFINITION 4.6. *Un lien par est dit scindant si et seulement si les deux arêtes qui le constituent séparent ses prémisses de sa conclusion.*



LEMME 4.7. *Si un réseau Π de taille l contient un lien par scindant et si tout réseau Π' de taille $l' < l$ est séquentialisable, alors Π est séquentialisable.*

DÉMONSTRATION: On suppose que Π contient un par scindant $A \bowtie B$. La suppression des deux arêtes de ce lien donne deux réseaux plus petits: l'un Π_1 de conclusions $C_1, \dots, C_k, A, B$ et d'hypothèses $H_1, \dots, H_u$, l'autre Π_2 de conclusions $C_{k+1}, \dots, C_n$ et d'hypothèses $H_{u+1}, \dots, H_p, A \bowtie B$ — en renumérotant éventuellement les C_i et les H_i . Comme Π_1 et Π_2 sont de taille plus petite que l (le lien par scindant a été supprimé), on dispose de deux preuves séquentielles π_1 et π_2 correspondant respectivement à Π_1 et Π_2 . La preuve π_2 utilise donc une hypothèse $\vdash A \bowtie B$. On obtient une preuve correspondant à Π de la sorte: on applique la règle PAR à

π_1 et on obtient une preuve π'_1 de $C_1, \dots, C_k, A \wp B$ sous les hypothèses $H_1, \dots, H_u$. On place π'_1 au dessus de l'hypothèse $\vdash A \wp B$ de π_2 dans laquelle on remplace la formule $A \wp B$ par la séquence $C_1, \dots, C_k, A \wp B$. On obtient ainsi une preuve π'_2 de conclusions $C_1, \dots, C_k, C_{k+1}, \dots, C_n$ sous les hypothèses $\vdash H_1 \dots \vdash H_u \vdash H_{u+1} \dots \vdash H_p$. $\diamond$

LEMME 4.8. *Soit Π un réseau de taille l tel que tout réseau de taille $l' < l$ soit séquentialisable. Si Π possède un lien par final alors Π est séquentialisable.*

DÉMONSTRATION: On considère le réseau Π' plus petit en supprimant un lien par final de Π . En appliquant la règle **par** à une séquentialisation de Π' , on obtient une séquentialisation de Π . $\diamond$

LEMME 4.9. *Si un réseau Π de taille l contient un lien tenseur ou coupure scindant et si tout réseau Π' de taille $l' < l$ est séquentialisable, alors Π est séquentialisable.*

DÉMONSTRATION: Si Π contient un lien tenseur ou coupure scindant, on obtient, en le supprimant, deux réseaux plus petits Π_1 et Π_2 , qui ont, par hypothèse, des séquentialisations π_1 et π_2 . On obtient une séquentialisation de Π en appliquant la règle **TENSEUR** ou **COUPURE** du calcul des séquents à π_1 et π_2 . $\diamond$

LEMME 4.10. *Soit Π un réseau sans lien par. Alors Π est séquentialisable.*

DÉMONSTRATION: On procède par récurrence sur le nombre de liens d'un tel réseau. S'il n'y a aucun lien tenseur ou coupure, le réseau est réduit à une hypothèse ou un axiome; sa séquentialisation est évidente. Sinon, considérons un lien tenseur ou coupure final; celui-ci étant scindant, l'argument de la démonstration du lemme 4.9. permet de conclure. $\diamond$

On voit alors que n'importe lequel des deux lemmes suivants suffit à achever la démonstration du théorème de séquentialisation.

LEMME 4.11. *Tout réseau contenant au moins un lien par contient un lien par scindant.*

LEMME 4.12. *Tout réseau dont toutes les conclusions sont des tenseurs contient un lien tenseur scindant.*

Nous allons les démontrer tous deux dans le paragraphe suivant. On remarquera qu'ils fournissent des algorithmes différents pour séquentialiser un réseau.

La notion de **par** scindant, et celle de **bloc** qui va suivre sont due à Vincent Danos. On trouvera aussi dans sa thèse [Dan90] une preuve indépendante de l'existence d'un **par** scindant. La nouveauté de notre travail sur ce sujet est de démontrer directement l'existence d'un tenseur scindant, de faire de ces deux techniques de séquentialisation qui jusque là semblaient sans rapport les corollaires d'un résultat général de combinatoire (celui du chapitre deux) et surtout d'affiner notre connaissance des chemins bigarrés (ou réalisables) dans un réseau avec **MÉLANGE**, ce qui aide à l'étude des réseaux ordonnés. Signalons aussi l'existence d'une autre preuve de la séquentialisation de ces réseaux, la première qui ait été trouvée, qui n'utilise pas cette notion, et ne semble pas être un corollaire du théorème du chapitre 2, due à moi-même et Arnaud Fleury [FR90].

§B. Réseaux et Agrégats

On va définir deux manières de transformer un réseau en un agrégat sans cycle bigarré, correspondant aux deux lemmes sus-mentionnés. Le résultat sera à chaque fois fourni par le théorème 2.4. du chapitre 2. La traduction donne une surjection s des sommets du réseau sur les sommets de l'agrégat, et la seconde une bijection des axiomes dans les sommets de l'agrégat (ce qui est aussi une surjection des feuilles du réseau sur les sommets de l'agrégat) ; ces surjections sont fidèles dans le sens suivant: l'agrégat contient un chemin bigarré entre $s(X)$ et $s(Y)$ si et seulement si le réseau contient un chemin bigarré entre X et Y .

¶B.1. Existence d'un par Scindant

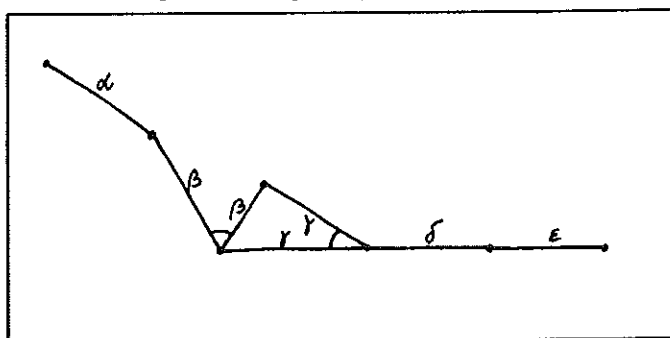
DÉFINITION 4.13. Un bloc d'un réseau est un de ses sous-graphes noirs connexes maximaux. On note $\mathcal{N}$ la surjection de l'ensemble des sommets de Π sur l'ensemble des blocs de Π qui à un sommet U de Π associe le bloc $\mathcal{N}(U)$ le contenant.

DÉFINITION 4.14. Soit Π un réseau, dont les couleurs propres sont $\alpha_1, \dots, \alpha_n$. On appelle traduction de Π suivant ses liens par γ , que l'on note Π_γ le graphe coloré suivant:

- * les sommets de Π_γ sont les blocs de Π .
- * les couleurs de Π_γ sont les couleurs propres de Π
- * on a une α -arête entre deux blocs X et Y de Π si et seulement si Π contient une α -arête entre un sommet de X et un sommet de Y .

En d'autres termes on obtient Π_γ à partir de Π en contractant ses blocs.

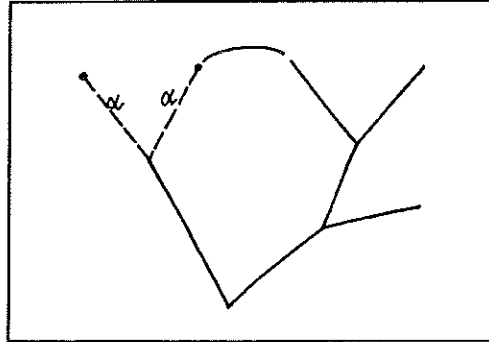
Donnons la traduction de l'exemple du chapitre précédent:



PROPOSITION 4.15. Soit Π un réseau; alors Π_γ est un agrégat de $K^{1,2}$ et de $K^{1,1}$.

DÉMONSTRATION: Soit α une couleur du réseau Π , correspondant au lien de prémisses U et V et de conclusion $W = U \wp V$. Il est clair que le graphe Π_γ contient au plus deux α -arêtes, les traces des deux α -arêtes de Π , et que celles-ci sont adjacentes (elles l'étaient dans Π , et Π_γ est une contraction de Π). On notera qu'il est cependant possible que Π_γ contienne une seule α arête — si U et V sont dans le même bloc, i.e si $\mathcal{N}(U) = \mathcal{N}(V)$.

Montrons maintenant que $\Pi_{\mathbf{p}}$ ne contient pas de boucle. Si $\Pi_{\mathbf{p}}$ contenait la boucle XX , c'est qu'une α -arête de Π avait ses deux sommets U et W dans $X = \mathcal{N}(U) = \mathcal{N}(W)$. Comme Π est un graphe simple, on sait que $U \neq W$, et comme U et W sont dans le même bloc X de Π , il existe un chemin noir p entre U et V . Le cycle $UpVaU$ est un cycle bigarré de Π , ce qui contredit l'hypothèse Π est un réseau.



◇

LEMME 4.16. Les propriétés suivantes sont équivalentes:

- (i) deux sommets X et Y de $\Pi_{\mathbf{p}}$ sont reliés par un chemin bigarré de $\Pi_{\mathbf{p}}$
- (ii) un sommet de X (i.e. un sommet U de Π tel que $\mathcal{N}(U) = X$) est relié par un chemin bigarré de Π à un sommet de Y (i.e. un sommet V de Π tel que $\mathcal{N}(V) = Y$).
- (iii) tout sommet de X est relié par un chemin bigarré de Π à un sommet de Y .

DÉMONSTRATION:

(ii) $\Rightarrow$ (iii) Soient U tel que $\mathcal{N}(U) = X$, V tel que $\mathcal{N}(V) = Y$ et p un chemin bigarré de Π reliant U et V . Soient U' tel que $\mathcal{N}(U') = \mathcal{N}(U) = X$, V' tel que $\mathcal{N}(V') = \mathcal{N}(V) = Y$. Comme U et U' (resp. V et V') sont dans le même bloc, il existe un chemin y (resp. z) noir reliant U et U' (resp. V et V'). Le chemin $U'yUpVzV'$ est bigarré (les arêtes ajoutées au chemin bigarré p sont noires).

(i) $\Rightarrow$ (ii) Soit $p = X_1 \dots X_n$ un chemin bigarré de $\Pi_{\mathbf{p}}$. On sait que deux sommets X_i et X_{i+1} de $\Pi_{\mathbf{p}}$ sont reliés par une α_i -arête si et seulement s'il existe deux sommets U_i et V_{i+1} de Π reliés par une α_i -arête a_i de Π tels que $\mathcal{N}(U_i) = X_i$ et $\mathcal{N}(V_{i+1}) = X_{i+1}$. Comme $\mathcal{N}(U_{i+1}) = \mathcal{N}(V_i + 1) = X_{i+1}$, les deux sommets U_i et V_i de Π sont reliés par un chemin noir w_i de Π . Le chemin $U_1, a_1, V_1, w_1, U_2, a_2, V_2, w_2, \dots, V_{n-1}, w_{n-1}, U_n, a_n, V_n$ est clairement un chemin bigarré de Π reliant le sommet U_1 de X_1 au sommet V_n de X_n .

(ii) $\Rightarrow$ (i) Soit $p = U_1 \dots U_n$ un chemin bigarré de Π reliant U_1 et U_n . On obtient un chemin bigarré reliant $\mathcal{N}(U_1)$ et $\mathcal{N}(U_n)$ en considérant le chemin $p' = \mathcal{N}(U_1) \dots \mathcal{N}(U_n)$ où les arêtes noires sont remplacées par le signe "=". ◇

PROPOSITION 4.17. Si Π est un réseau, alors $\Pi_{\mathbf{p}}$ ne contient pas de cycle bigarré.

DÉMONSTRATION: Conséquence directe du lemme précédent. ◇

PROPOSITION 4.18. Soit Π un réseau et Π_v sa traduction suivant ses liens par $\cdot$. Un lien par de couleur α est scindant si et seulement si la couleur α est une couleur scindante de l'agrégat Π_v .

DÉMONSTRATION: S'il existe un chemin (par forcément bigarré) entre les deux cotés d'un des $K^{1,2}$ (ou $K^{1,1}$) de Π_v , alors ce chemin correspond à un chemin de Π — en remplaçant les points de Π_v par des chemins noirs. De même, un chemin de Π correspond à un chemin de Π_v , en remplaçant les arêtes noires par des signes "=". $\diamond$

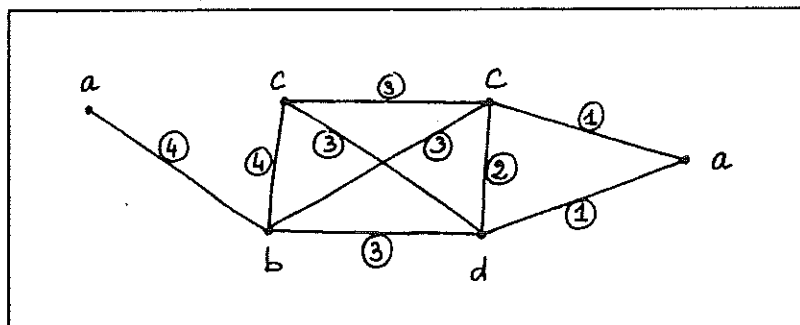
DÉMONSTRATION DU LEMME 4.11. En appliquant le théorème 2.4. à une composante connexe de Π_v non réduite à un sommet, on obtient un lien par scindant dans l'une des composantes connexes de Π , qui est bien évidemment un lien par scindant de Π .

¶B.2. Existence d'un tenseur Scindant

La deuxième traduction envisagée donne un agrégat dont les sommets sont les arêtes axiome du réseau, et les couleurs les liens tenseurs du réseau.

DÉFINITION 4.19. Soit Π un réseau, soit Ax l'ensemble de ses arêtes axiomes et soit $\mathcal{L}_\otimes$ l'ensemble de ses liens tenseur dont les éléments seront notés $\otimes_\alpha, \otimes_\beta \dots$ etc. On définit alors le graphe $\Pi_\otimes$ appelé traduction de Π suivant ses liens tenseur comme un graphe coloré dont les sommets sont Ax et les couleurs $\mathcal{L}_\otimes$. On a une arête de couleur $\otimes_\alpha \in \mathcal{L}_\otimes$ entre deux de ses sommets $a \in Ax$ et $a' \in Ax$ si et seulement si l'un des deux sommets de l'arête a est au dessus de l'une des prémisses du lien tenseur $\otimes_\alpha$ de Π et si l'un des deux sommets de a' est au dessus de l'autre prémisses du lien tenseur $\otimes_\alpha$ de Π .

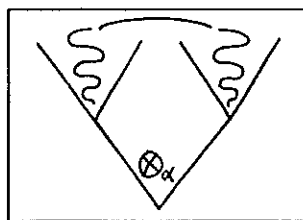
Donnons la traduction de l'exemple du chapitre précédent:



PROPOSITION 4.20. Le graphe coloré $\Pi_\otimes$ est un agrégat.

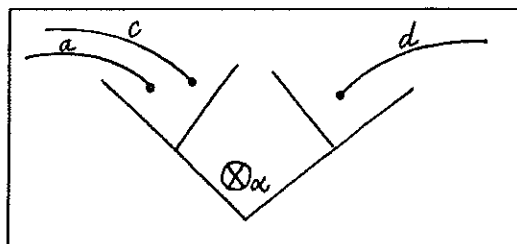
DÉMONSTRATION: On montre d'abord qu'il n'y a pas de boucle dans $\Pi_\otimes$. Ensuite, il suffit de montrer que, pour toute couleur $\otimes_\alpha$ de $\Pi_\otimes$ et pour tout couple d' $\otimes_\alpha$ -arêtes ab et cd de $\Pi_\otimes$, ac ou ad est une $\otimes_\alpha$ -arête de $\Pi_\otimes$, et que, pour tout couple d' $\otimes_\alpha$ -arêtes ab et ac de $\Pi_\otimes$, bc n'est pas une $\otimes_\alpha$ -arête de $\Pi_\otimes$.

$\Pi_\otimes$ ne contient pas de boucle Si $\Pi_\otimes$ contenait la boucle aa de couleur $\otimes_\alpha$, alors l'arête axiome a aurait un de ses sommets au dessus de l'une des prémisses du lien tenseur $\otimes_\alpha$ et l'autre au dessus de l'autre prémisses de ce lien. Il y aurait alors un cycle bigarré dans le réseau, vu que les deux sous-arbres au dessus des prémisses de ce lien ne contiennent pas deux fois la même couleur, et que les arêtes du lien tenseur sont noires.

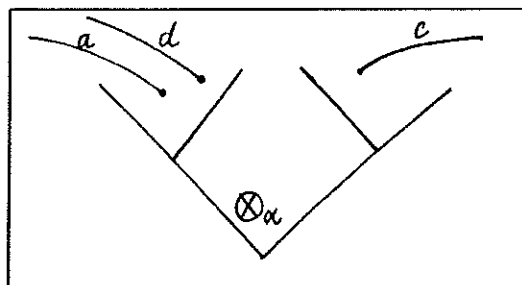


$ab, cd \in \otimes_\alpha \Rightarrow (ac \in \otimes_\alpha \vee ad \in \otimes_\alpha)$ Supposons que $ab, cd \in \otimes_\alpha$. Il existe donc un sommet A de l'arête axiome a du réseau et un sommet B de l'arête axiome b du réseau qui soient l'un au dessus d'une des prémisses du lien tenseur $\otimes_\alpha$, et l'autre au dessus de l'autre prémissse du lien tenseur $\otimes_\alpha$. Symétriquement, il existe un sommet C de l'arête axiome c du réseau et un sommet D de l'arête axiome d du réseau qui soient l'un au dessus d'une des prémisses du lien tenseur $\otimes_\alpha$, et l'autre au dessus de l'autre prémissse du lien tenseur $\otimes_\alpha$. Un lien tenseur n'ayant que deux prémisses on voit que:

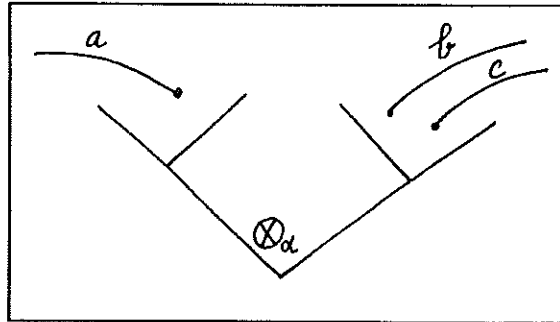
★ ou bien A et C sont au dessus de la même prémissse du lien tenseur $\otimes_\alpha$ auquel cas A et D sont au dessus de prémisses différentes, et on a une $\otimes_\alpha$ -arête ad dans $\Pi_\otimes$.



★ ou bien A et D sont au dessus de la même prémissse du lien tenseur $\otimes_\alpha$ auquel cas A et C sont au dessus de prémisses différentes, et on a une $\otimes_\alpha$ -arête ac dans $\Pi_\otimes$.



$ab, ac \in \otimes_\alpha \Rightarrow bc \notin \otimes_\alpha$ Supposons que $ab, ac \in \otimes_\alpha$. Il existe donc un sommet A de l'arête axiome a du réseau et un sommet B de l'arête axiome b du réseau qui soient l'un au dessus d'une des prémisses du lien tenseur $\otimes_\alpha$ et l'autre au dessus de l'autre prémissse du lien tenseur $\otimes_\alpha$. Symétriquement, il existe un sommet A' de l'arête axiome a du réseau et un sommet C de l'arête axiome c du réseau qui soient l'un au dessus d'une des prémisses du lien tenseur $\otimes_\alpha$ et l'autre au dessus de l'autre prémissse du lien tenseur $\otimes_\alpha$. Comme A et A' ne sont pas au dessus des prémisses différentes du lien tenseur $\otimes_\alpha$ — sinon on aurait une $\otimes_\alpha$ -boucle aa , ce qu'on a réfuté précédemment — B et C sont au dessus de la même prémissse du lien tenseur $\otimes_\alpha$. Si l'on avait une $\otimes_\alpha$ arête bc dans $\Pi_\otimes$, b ou c aurait son autre extrémité au dessus de l'autre prémissse de ce lien tenseur $\otimes_\alpha$, et on a déjà vu que c'est impossible.



◇

PROPOSITION 4.21. *Deux axiomes a et a' de Π sont connectés dans Π si et seulement si les sommets correspondants de $\Pi_\otimes$ le sont dans $\Pi_\otimes$.*

DÉMONSTRATION: S'ils sont connectés dans $\Pi_\otimes$, il est clair qu'ils le sont dans Π , chaque $\otimes_\alpha$ arête de $\Pi_\otimes$ se mimant par un chemin de Π .

Supposons qu'on ait un chemin p entre a et a' dans Π . On procède par induction sur le nombre d'axiomes rencontrés entre a et a' . Si c'est 0, le chemin entre a et a' est tout entier contenu dans un arbre et ce chemin est donc unique. Si son point le plus bas est un lien tenseur $\otimes_\beta$, les deux sommets a et a' de $\Pi_\otimes$ sont adjacents par une $\otimes_\beta$ -arête. Si le point le plus bas est un lien *par*, il existe un lien tenseur en dessous, disons $\otimes_\gamma$. Soit g un axiome situé au dessus de l'autre coté de $\otimes_\gamma$; alors $\Pi_\otimes$ contient les $\otimes_\gamma$ -arêtes ga et ga' ; les sommets a et a' de $\Pi_\otimes$ sont donc connectés dans $\Pi_\otimes$. La connexité étant transitive, une récurrence triviale achève la démonstration. ◇

PROPOSITION 4.22. *Les deux prémisses A et B d'un lien tenseur $\otimes_\alpha$ sont connectés dans Π par un chemin n'utilisant aucune de ses deux arêtes si et seulement s'il existe deux axiomes a et b , a au dessus de A et b au dessus de B tels que les sommets correspondants de $\Pi_\otimes$ soient connectés dans $\Pi_\otimes$ par un chemin n'utilisant aucune arête de couleur $\alpha_\otimes$.*

DÉMONSTRATION: Supposons que l'on ait un chemin entre A et B dans Π n'utilisant aucune des deux arêtes du lien $\otimes_\alpha$. En raccourcissant éventuellement ce chemin, on peut supposer qu'il existe dans Π un chemin entre deux axiomes a et a' situés de part et d'autre de $\alpha_\otimes$ ne passant pas par $\otimes_\alpha$. Considérons le réseau Π' obtenu à partir de Π en supprimant le lien tenseur $\otimes_\alpha$ et les liens *par* situés au dessus de lui. Les axiomes a et a' sont connectés dans ce réseau dont toutes les conclusions sont des tenseurs, et la proposition précédente montre qu'ils le sont dans $\Pi'_\otimes$; comme $\Pi'_\otimes$ contient $\Pi_\otimes$ (mêmes sommets mais plus d'arêtes) et que $\otimes_\alpha$ n'est pas une couleur de $\Pi'_\otimes$, les points a et a' sont connectés dans $\Pi_\otimes$ par un chemin n'utilisant pas d' $\otimes_\alpha$ -arêtes. ◇

PROPOSITION 4.23. (i) *Soit p un chemin bigarré de $\Pi_\otimes$ reliant ses sommets a et a' ; alors Π contient un chemin bigarré reliant un sommet de l'axiome a et un sommet de l'axiome a' .*

(ii) *Soit p un chemin bigarré de Π reliant un sommet de l'axiome a et un sommet de l'axiome a' ; alors $\Pi_\otimes$ contient un chemin bigarré reliant ses sommets a et a' .*

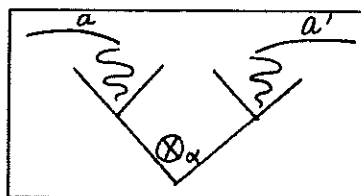
Pour démontrer cette proposition, nous avons besoin du lemme suivant:

LEMME 4.24. *Soient x et y deux feuilles situées au dessus d'un tenseur $\otimes_\alpha$ dans Π ; s'il existe un chemin bigarré entre x et y ne passant pas par $\otimes_\alpha$, alors x et y sont d'un même côté du tenseur $\otimes_\beta$.*

DÉMONSTRATION: On suppose qu'il existe un couple de feuilles situées de part et d'autre du tenseur $\otimes_\beta$ et on en choisit un x, y de sorte que le chemin entre x et y ne passant pas par $\otimes_\alpha$ contienne un nombre minimal de feuilles de l'arbre A_β . Si ce chemin ne passe pas au dessus $\otimes_\beta$, alors il y a un cycle bigarré. S'il passe au dessus de β , alors il contient un z au dessus de $\otimes_\beta$; ce z est soit du même côté de $\otimes_\beta$ que x , soit du même côté de $\otimes_\beta$ que y . Les feuilles x, z (ou z, y) sont situées de part et d'autre de $\otimes_\beta$ et il existe un chemin les reliant contenant strictement moins de feuilles de A_β , ce qui contredit l'hypothèse de minimalité. $\diamond$

DÉMONSTRATION:

- (i) On procède par induction sur la longueur l du chemin bigarré reliant a et a' dans $\Pi_\otimes$ pour établir qu'il existe un tel chemin où toute $\otimes_\alpha$ -arête est remplacée par un V-chemin de rebond $\otimes_\alpha$. Si $l = 1$ le résultat est évident, vu que les deux arbres premisses du tenseur $\otimes_\alpha$ n'ont pas de couleur commune, et que les deux chemins reliant la conclusion du tenseur $\otimes_\alpha$ à l'axiome a et à l'axiome a' sont bigarrés.

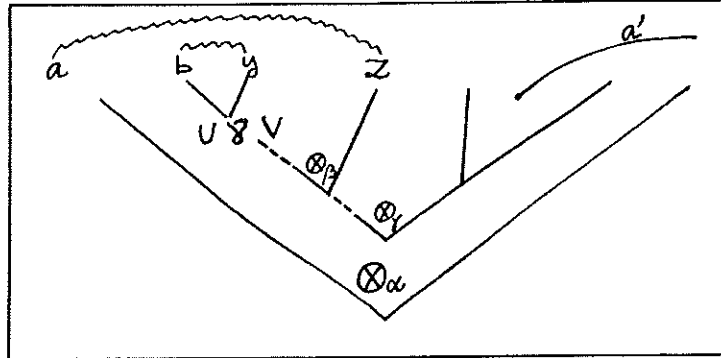


Soit ba' de couleur $\otimes_\beta$, la dernière arête de p dans $\Pi_\otimes$. D'après l'hypothèse d'induction, on dispose d'un chemin bigarré s de Π , de l'axiome a à l'axiome b . Appelons $\otimes_\alpha$ le tenseur final situé en dessous de $\otimes_\beta$ et v le chemin bigarré de A_α reliant b et a' et passant par $\otimes_\beta$. Si s ne passe pas par A_α , alors le prolongement de asb par v de b à a' contenu dans A_α est un chemin bigarré de Π .

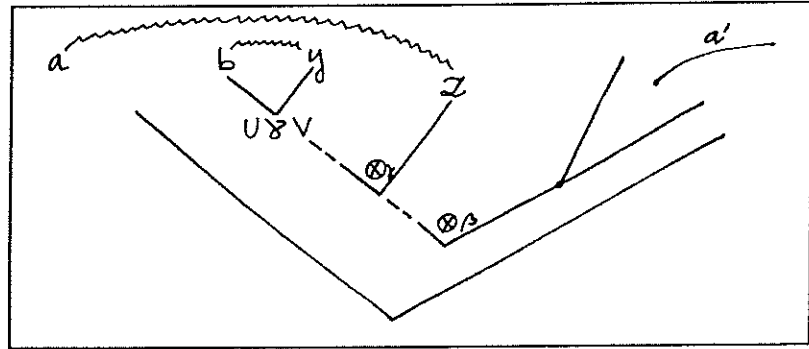
S'ils n'empruntent pas d'arêtes de Π de même couleur, $asbva'$ convient.

Sinon, c'est que s et t contiennent un point en commun, et soit $U\mathfrak{V}V$ leur premier point commun sur s qui est nécessairement la conclusion d'un par de A_α . Appelons x la première feuille de A_α qui précède $U\mathfrak{V}V$, et y la première feuille de A_α qui suit $U\mathfrak{V}V$; on a alors $as'xs''ys'''b$. L'unique chemin de A_α entre x et y passe par un tenseur $\otimes_\gamma$. Il s'ensuit que x et y sont de part et d'autre de $\otimes_\gamma$, et que b et a' sont de part et d'autre du tenseur $\otimes_\beta$. Notons aussi qu'en vertu de l'existence du chemin bigarré entre y et b ne passant pas par $\otimes_\beta$, y et b sont d'un même côté de $\otimes_\beta$. Pour la même raison, y et b sont aussi d'un même côté de $\otimes_\gamma$. Comme le par $U\mathfrak{V}V$ est au dessus et de $\otimes_\beta$ et de $\otimes_\gamma$, l'un des deux tenseur $\otimes_\beta$ et $\otimes_\gamma$ est au dessus de l'autre.

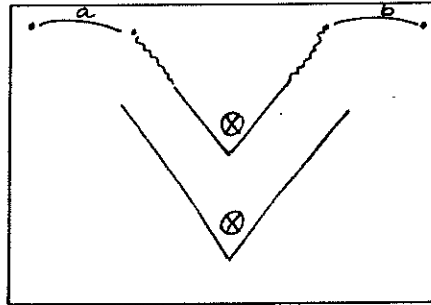
Si $\otimes_\gamma$ est au dessous de $\otimes_\beta$, comme y et x sont de part et d'autre de $\otimes_\gamma$ et y et b de part et d'autre de $\otimes_\beta$, $\otimes_\beta$ est au dessus de $\otimes_\gamma$ du côté de y ; donc a' et x sont de part et d'autre de $\otimes_\gamma$. Le chemin cherché est $as'x(\otimes_\gamma)a'$.



Si $\otimes_\gamma$ est au dessus de $\otimes_\beta$, on montre de même que x et a' sont de part et d'autre de $\otimes_\beta$. Le chemin cherché est $as'x(\otimes_\beta)a'$



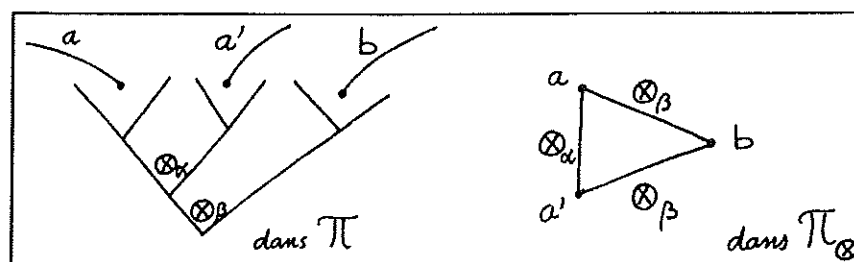
- (ii) Il suffit de remarquer que toute partie maximale du chemin contenue dans un même arbre peut être remplacée par une arête de $\Pi_\otimes$, et qu'on obtient alors un chemin bigarré de Π .



◇

LEMME 4.25. Soit Π un réseau. Un lien tenseur $\otimes_\alpha$ d'un réseau Π est scindant si et seulement si cette couleur $\otimes_\alpha$ est scindante dans l'agrégat correspondant $\Pi_\otimes$.

DÉMONSTRATION: Montrons, pour commencer, que si le tenseur $\otimes_\alpha$ du réseau Π n'est pas final, la couleur $\otimes_\alpha$ n'est pas scindante dans l'agrégat $\Pi_\otimes$. Si le lien tenseur $\otimes_\alpha$ n'est pas final, alors il est au dessus d'un lien tenseur final $\otimes_\beta$. Soient b un axiome dont une conclusion est au dessus de l'autre prémisses du lien tenseur $\otimes_\beta$, a un axiome dont une conclusion est au dessus de la prémisses gauche du tenseur $\otimes_\alpha$, et a' un axiome dont une conclusion est au dessus de la prémisses droite du tenseur $\otimes_\alpha$. Il s'ensuit que dans $\Pi_\otimes$, les arêtes ab et $a'b$ sont de couleur $\otimes_\beta$ et que l'arête aa' est de couleur $\otimes_\alpha$, ce qui montre que la couleur $\otimes_\alpha$ n'est pas scindante dans $\Pi_\otimes$.



Supposons maintenant que le tenseur $\otimes_\alpha$ soit final mais pas scindant dans Π . Il existe donc dans Π un chemin (nécessairement non bigarré) d'une arête axiome a dont un sommet est au dessus d'une des prémisses de ce lien tenseur à une arête axiome a' dont un sommet est au dessus de l'autre prémisses de ce lien tenseur, n'empruntant pas les arêtes de ce lien tenseur. Ce chemin correspond (d'après la proposition précédente) à un chemin dans $\Pi_\otimes$ n'utilisant pas d' $\otimes_\alpha$ -arête entre a et a' , et comme aa' est une $\otimes_\alpha$ -arête de $\Pi_\otimes$, cela suffit à assurer que $\otimes_\alpha$ n'est pas une couleur scindante de $\Pi_\otimes$.

Montrons maintenant que si une couleur de $\Pi_\otimes$ correspond à un lien tenseur scindant de Π alors cette couleur est une couleur scindante de $\Pi_\otimes$. S'il existait dans $\Pi_\otimes$ un chemin n'utilisant pas d' $\otimes_\alpha$ -arête entre les deux axiomes a et a' , a dont l'un des sommets soit au dessus d'une partie de $\otimes_\alpha$ et a' au dessus de l'autre prémisses du lien tenseur $\otimes_\alpha$ alors, d'après les propositions précédentes il existerait dans Π un chemin n'utilisant pas de $\otimes_\alpha$ -arête entre a et a' . $\diamond$

DÉMONSTRATION DU LEMME 4.12. D'après le lemme précédent, il suffit d'appliquer le théorème 2.4..

¶B.3. Optimisation de la méthode du tenseur scindant

PROPOSITION 4.26. Soit Π un réseau dont toutes les conclusions soient des tenseurs. Considérons Π' le réseau obtenu en remplaçant tout lien tenseur non final par un lien par. Alors Π' est un réseau et de plus $\otimes_\alpha$ est un tenseur scindant de Π' équivaut à $\otimes_\alpha$ est un tenseur scindant de Π .

DÉMONSTRATION: Les chemins bigarrés de Π contiennent ceux de Π' , et Π' ne contient donc aucun cycle bigarré. Les chemins (pas forcément bigarrés) de Π et Π' étant les mêmes une arête est un isthme de Π si et seulement si c'est un isthme de Π' , ce qui achève la démonstration. $\diamond$

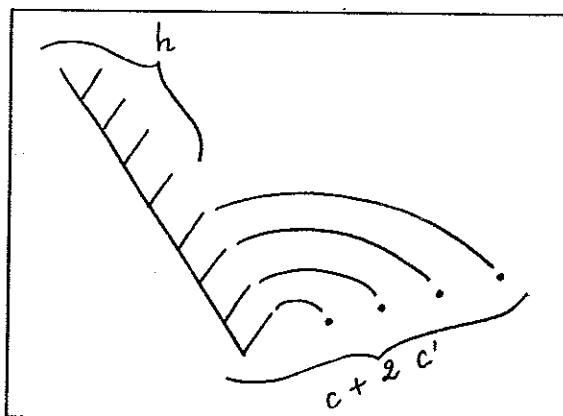
L'algorithme pour trouver un tenseur scindant fonctionne donc en un temps $\frac{t^3}{6}$ où t est le nombre de tenseur finaux. Il a été dit que l'algorithme pour trouver un tenseur scindant était linéaire car il se ramenait à trouver un isthme dans un graphe: c'est faux, car si l'isthme trouvé est un axiome, ou même une arête d'un lien non-final, cet isthme ne permet pas de séquentialiser le réseau, en tout cas pas directement.

¶B.4. Agrégats et réseaux

PROPOSITION 4.27. Soit $\mathcal{G}$ un agrégat sans cycle bigarré; alors il existe un réseau Π tel que $\Pi_\nabla = \mathcal{G}$

DÉMONSTRATION: Un sommet de $\mathcal{G}$ est dit prémisses d'un $K^{1,2}$ si et seulement s'il est l'un des deux points de sa partie à deux éléments, et il est dit conclusion s'il est le sommet de sa

partie à un élément. Pour chaque $K^{1,1}$ de $\mathcal{G}$, appelons (arbitrairement) prémisses l'un des deux sommets et conclusion l'autre. Soit maintenant X un sommet de $\mathcal{G}$, qui est c fois prémisses d'un $K^{1,2}$, c' fois d'un $K^{1,1}$, et h fois conclusion d'un $K^{1,2}$ ou $K^{1,1}$. On considère le bloc suivant:



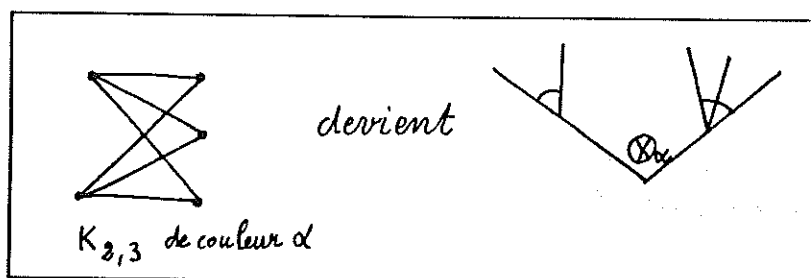
Il est clair qu'en plaçant les par entre les points correspondants de blocs ainsi construits pour tout sommet de $\mathcal{G}$, on obtient un réseau dont la traduction suivant les par est l'agrégat de départ. $\diamond$

La proposition suivante montre que tout agrégat sans cycle bigarré est grosso modo la traduction selon les tenseurs d'un réseau. On peut vraisemblablement montrer par un argument plus fin qu'il est exactement la traduction selon les liens tenseurs d'un réseau. Nous pensons que cela n'en vaut pas la peine, l'essentiel étant de se convaincre que réseaux et agrégats sont combinatoirement les mêmes objets, ce que la proposition suivante suffit à établir.

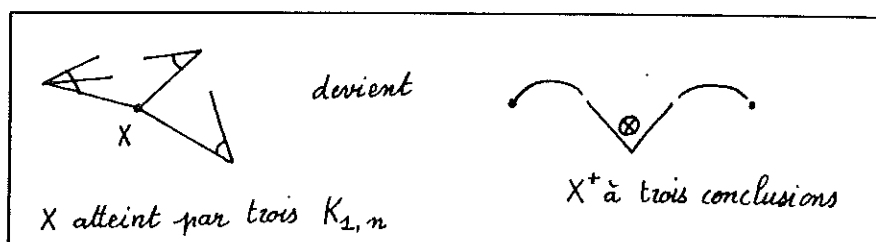
PROPOSITION 4.28. Soit $\mathcal{G}$ un agrégat sans cycle bigarré; alors il existe un réseau Π et une surjection $(...)^+$ de $\Pi_{\otimes}$ dans $\mathcal{G}$ tels que:

X et X' sont reliés par un chemin bigarré dans $\mathcal{G}$ si et seulement si tout axiome de X^+ et tout axiome de X'^+ sont reliés par un chemin bigarré dans $\Pi_{\otimes}$. (cela signifie que $\Pi_{\otimes}$ est essentiellement semblable à $\mathcal{G}$)

DÉMONSTRATION: On commence par transformer l'agrégat ainsi: pour tout $K^{n,p}$ de l'agrégat on ajoute deux arêtes entre les deux parties du $K^{n,p}$, ce qui correspond à un tenseur entre deux par n -aires.



Puis on remplace tout point X atteint par p $K^{1,n}$ par : le bloc X^+ suivant:



On obtient alors un réseau Π dont l'agrégat $\Pi_{\otimes}$ n'est pas tout à fait $\mathcal{G}$ mais qui satisfait: X et X' sont reliés par un chemin bigarré dans $\mathcal{G}$ si et seulement si tout axiome de X^+ et tout axiome de X'^+ sont reliés par un chemin bigarré dans $\Pi_{\otimes}$. $\diamond$

¶B.5. Existence d'un tenseur héréditairement scindant

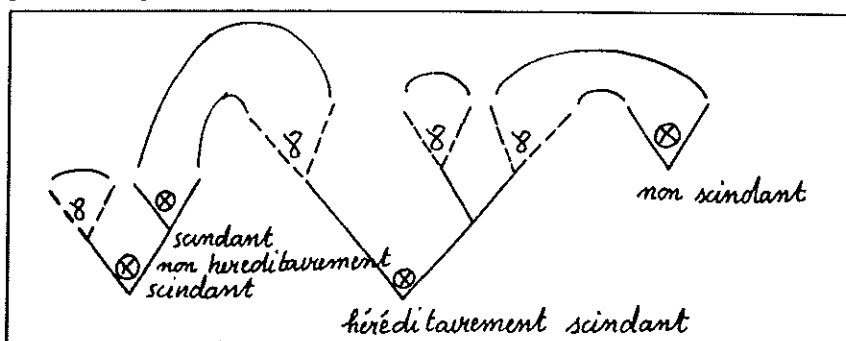
Nous établissons dans ce paragraphe un résultat nécessaire à la démonstration du théorème 7.24. du paragraphe §F. où nous explicitons le rapport entre ce calcul-ci et le calcul ordonné.

DÉFINITION 4.29. On dit qu'un lien tenseur (ou coupure) t d'un réseau Π est héréditairement scindant si et seulement s'il satisfait la condition récursive suivante:

- * t est scindant dans Π
- * si t' est un lien tenseur dont la conclusion est une prémisse de t , alors t' est héréditairement scindant dans la composante connexe du réseau $\Pi - t$ le contenant — où $\Pi - t$ désigne le réseau obtenu en supprimant le lien tenseur final t .

Il est alors clair que "le lien tenseur t est héréditairement scindant dans sa composante connexe" équivaut à "le lien tenseur t est héréditairement scindant".

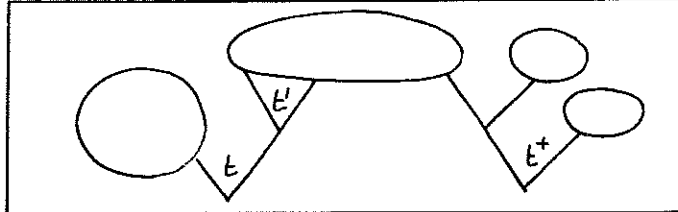
Donnons un petit exemple:



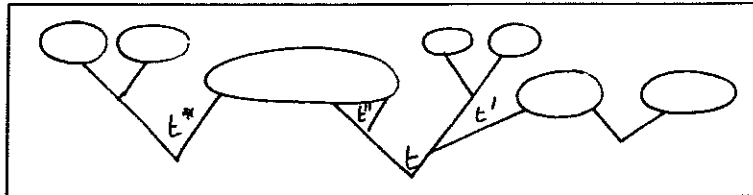
LEMME 4.30. [du tenseur héréditairement scindant] Dans un réseau dont toutes les conclusions sont des tenseur il existe un lien tenseur héréditairement scindant.

DÉMONSTRATION: On procède par induction sur le nombre de liens tenseur du réseau. Si c'est un, ce lien tenseur est final, scindant et n'a pas de prémisse qui soit un lien tenseur: il est donc héréditairement scindant. Sinon, soit t un tenseur scindant de Π ; si aucune des deux prémisses de t n'est un tenseur, alors t est héréditairement scindant dans Π . Sinon, soit t' un lien tenseur prémisse de t ; considérons alors Π' la composante connexe de t' dans $\Pi - t$ (dont toutes les conclusions sont des tenseur), qui contient par hypothèse d'induction un

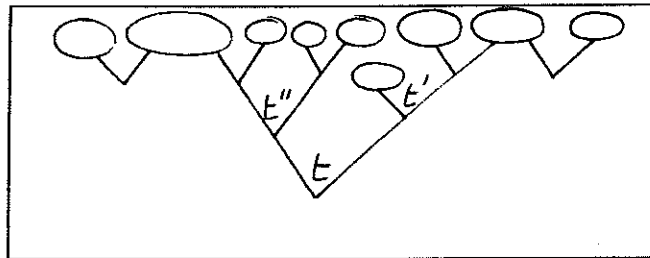
lien tenseur héréditairement scindant t^+ ; si $t^+ \neq t'$ alors t^+ est héréditairement scindant dans Π .



On peut donc supposer désormais que t' est héréditairement scindant dans Π' ; si l'autre prémisses de t n'est pas un tenseur, t est alors héréditairement scindant dans Π .



Sinon, soit t'' l'autre prémisses de t ; considérons alors Π'' la composante connexe de t'' dans $\Pi - t$ (dont toutes les conclusions sont des tenseurs), qui contient par hypothèse d'induction un lien tenseur héréditairement scindant t^* . Si $t^* \neq t''$ alors ce lien tenseur est héréditairement scindant dans Π . Si $t^* = t''$ le lien t'' est héréditairement scindant dans Π'' , et comme le lien t' l'est dans Π' , le lien t est héréditairement scindant dans Π .



◇

Chapitre 5

Éléments neutres

Nous allons montrer que la règle de MÉLANGE permet de construire un calcul en réseaux avec les éléments neutres du **par** ($\perp$) et du **tenseur** ($\otimes$) où il n'y ait pas de boîte $\perp$ et dans le quel ces deux neutres soient différents.

La solution que la règle de MÉLANGE suggère est pourtant d'élaborer un calcul où $\perp = 1$. C'est peu satisfaisant car cela signifie qu'une formule affaiblie se comporte comme une formule de connecteur principal bien-sûr ($!A$), ce qui entraîne notamment l'existence de preuves de conclusion $!A \wp !B$, dès que $!A$ et $!B$ sont prouvable, alors que le calcul usuel ne contient jamais de preuve une telle conclusion. Le calcul usuel a par contre l'inconvénient d'être trop complexe pour avoir un bon critère des réseaux puisque la prouvabilité des formules écrites exclusivement avec $\perp$, 1 , $\wp$ et $\otimes$ est un problème NP-complet [LW92].

Nous montrons ici que la fin du chapitre 2 (le théorème 2.28. permet une solution intermédiaire. Cette solution utilise un entier associé au réseau, appelé indice du réseau, qui est préservé par élimination des coupures: le nombre de $\perp$ moins le nombre de composantes connexes de tout sous-graphe bigarré maximal (cette définition est licite car le théorème 2.28. montre que tous les sous-graphes bigarrés maximaux ont le même nombre de composantes connexes). Le critère est alors le suivant:

(c) le réseau est sans circuit bigarré et son indice est 0.

Notons enfin que cet entier fournit de plus une sémantique des phases.

La portée de ce résultat dépasse le cadre multiplicatif dans le quel nous le présentons. En effet les boîtes $\perp$ sont évidemment les mêmes que les boîtes affaiblissements, et ces boîtes sont les seuls obstacles à la confluence du calcul. La généralisation de notre critère "sans cycle bigarré et nombre de" à un calcul avec $?, !, \forall, \exists$ *sans boîtes affaiblissement* (et donc confluent) est le critère récursif suivant. Un réseau est correct si et seulement si:

le remplacement de ces boîtes maximales par des axiomes n -aires satisfait (c);

l'intérieur de chaque boîte maximale est un réseau correct.

Notons que c'est là le seul critère connu pour reconnaître les réseaux des pré-réseaux (ou les "proofs" des "proof expressions" dans un calcul correspondant à [Abr91], i.e. à [Gir87a] sans boîtes affaiblissement. Le résultat négatif de [LW92] suggère qu'on ne peut obtenir de critère raisonnable décrivant exactement le calcul des séquents original sans introduire de boîtes affaiblissement et perdre la confluence. Ainsi, bien que nous ayons strictement plus de formules prouvables, (telle $\perp \otimes (1 \wp 1)$), notre solution semble optimale.

§A. Le Calcul des Séquents avec Mélange et Éléments Neutres

¶A.1. Une Sémantique des Phases

Nous renvoyons à [Gir87a] pour la définition précise d'un espace des phases. Disons simplement qu'on se donne un monoïde commutatif et une partie $\perp^\varphi$ de ce monoïde, l'ensemble des antiphasés. Deux éléments du monoïde sont dit orthogonaux si leur composé est une antiphasé, et on étend comme d'habitude cette notion aux sous-ensembles du monoïde. On interprète alors les formules par des faits, i.e. des sous-ensembles égaux à leur bi-orthogonal. Les formules vraies sont celles dont l'interprétation contient 1^φ , l'orthogonal des antiphasés.

Nous travaillerons ici avec des séquents indexés par un entier relatif. Cet entier peut se voir comme un fait d'une sémantique des phases fort naturelle due à M. Ajlani, que nous décrivons brièvement.

★ le monoïde est $(\mathbb{Z}, +)$ (addition usuelle sur les entiers relatifs)

★ l'ensemble des antiphasés $\perp^\varphi$ est $[1, +\infty)$

★ les faits sont les sections commençantes de $\mathbb{Z}$:

$$\begin{aligned} 1^\varphi = (\perp^\varphi)^\perp &= [0, +\infty) \\ [a, +\infty)^\perp &= [1 - a, +\infty) \\ [a, +\infty) \otimes [b, +\infty) &= [a + b, +\infty) \\ [a, +\infty) \wp [b, +\infty) &= [a + b - 1, +\infty) \end{aligned}$$

On identifiera n et $[n, +\infty)$.

¶A.2. Les Règles du Calcul des Séquents

Règles structurelles	
$\frac{\begin{array}{c} \vdots \pi_1 \\ \hline \vdots n \\ \vdots A_1, A_2, \dots, A_n \end{array} \quad \begin{array}{c} \vdots \pi_2 \\ \hline \vdots m \\ \vdots B_1, B_2, \dots, B_p \end{array}}{\vdots n+m-1 \\ \vdots A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p} \text{ MÉLANGE}$	
$\frac{}{\vdots 1} \text{ VIDE}$	

Règles d'identité	
$\frac{}{\vdots 0 \\ \vdots A, A^\perp} \text{ AXIOME}$	
$\frac{\begin{array}{c} \vdots \pi_1 \\ \hline \vdots n \\ \vdots A_1, A_2, \dots, A_n, K \end{array} \quad \begin{array}{c} \vdots \pi_2 \\ \hline \vdots m \\ \vdots K^\perp, B_1, B_2, \dots, B_p \end{array}}{\vdots n+m \\ \vdots A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p} \text{ COUPURE}$	

Règles logiques		
$\frac{\frac{\vdots \pi_1}{\vdash^n A_1, A_2, \dots, A_n, A} \quad \frac{\vdots \pi_2}{\vdash^m B, B_1, B_2, \dots, B_p}}{\vdash^{n+m} A_1, A_2, \dots, A_n, A \otimes B, B_1, B_2, \dots, B_p} \text{ TENSEUR}$		$\frac{}{\vdash^0 1} \text{ UN}$
$\frac{\frac{\vdots \pi_1}{\vdash^n A, B, A_1, A_2, \dots, A_n}}{\vdash^n A \wp B, A_1, A_2, \dots, A_n} \text{ PAR}$		$\frac{}{\vdash^1 \perp} \text{ FAUX}$

REMARQUE 5.1.

- * L'indice d'une preuve séquentielle π , noté $\#(\pi)$ est l'indice du séquent conclusion.
- * Ce calcul est évidemment valide pour l'interprétation dans le modèle précédemment décrit.
- * La règle habituelle du $\perp$ est une règle dérivée: MÉLANGE(π , FAUX), donne une preuve de même indice que π .
- * L'axiome du séquent vide VIDE est nécessaire à l'élimination des coupures: considérer, par exemple la preuve COUPURE(UN , FAUX)
- * Si l'on ommet les règles VIDE et MÉLANGE, mais ajoute la règle (dérivable) du $\perp$ l'indice est toujours nul, et on reconnaît le calcul ordinaire.
- * S'il on ommet les indices, on a $\perp \equiv 1$.
- * S'il on ommet les indices et les règles FAUX et UN, c'est le calcul des deux chapitres précédents.

DÉFINITION 5.2. On dit que F est valide au niveau i s'il existe une preuve de $\vdash^i F$.

F est valide signifie que F est valide au niveau 0.

REMARQUE 5.3. On a strictement plus de formules valides que de formules valides dans le calcul ordinaire: par exemple $\perp \otimes (1 \wp 1)$ et $((A \otimes B) \rightarrow (A \wp B)) \otimes \perp$ sont valides. On remarque de plus que $\perp$ et $(A \otimes B) \rightarrow (A \wp B)$ sont valides¹ au niveau 1.

THÉORÈME 5.4. (Élimination des coupures) Soit $\mathcal{P}$ une preuve séquentielle de $\vdash^n \Gamma$ alors il existe une preuve sans coupure de $\vdash^n \Gamma$.

DÉMONSTRATION: C'est un corollaire immédiat du théorème de séquentialisation présenté ultérieurement. On procède ainsi: on traduit la preuve séquentielle en son réseau associé, on normalise le réseau ainsi obtenu, et on lui associe (par le théorème de séquentialisation 5.16.) une preuve sans coupure. Nous montrerons dans les propositions 5.8, 5.10. et 5.16. que l'indice de la preuve est préservé par élimination des coupures. $\diamond$

¹De toute façon elles sont toutes interprétées par la sémantique cohérente

§B. Les Réseaux Correspondants

Réseaux et préréseaux sont les mêmes que ceux des deux chapitres précédents. Cependant un réseau sera dit correct si son indice, défini quelques lignes plus bas est 0. On aura de plus deux liens supplémentaires correspondants aux éléments neutres:

DÉFINITION 5.5. * On appelle lien 1 une hypothèse 1.

* On appelle lien $\perp$ une hypothèse $\perp$.

* On appelle $ff(\Pi)$ le nombre de liens $\perp$ d'un réseau Π

PROPOSITION 5.6. Soit Π un réseau. Le nombre de composantes connexes de ses sous-graphes bigarrés maximaux (ou graphes de correction) est toujours le même. On le notera $cc(\Pi)$.

DÉMONSTRATION: Conséquence directe de la proposition 2.28.. ◇

DÉFINITION 5.7. On appelle indice d'un réseau le nombre $\#(\Pi)$ suivant:

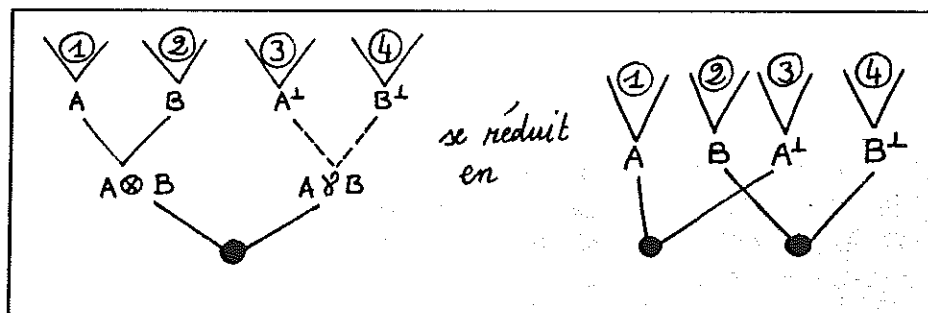
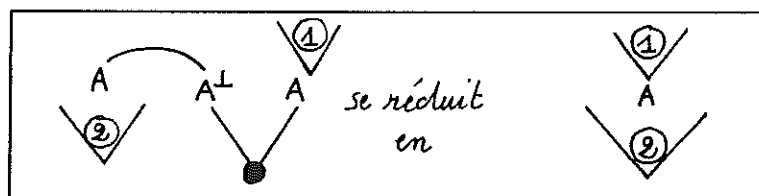
$$\#(\Pi) = ff(\Pi) - cc(\Pi) + 1$$

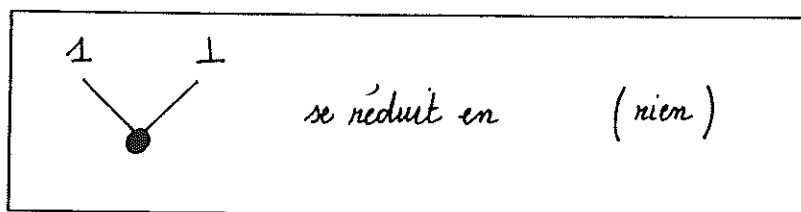
PROPOSITION 5.8. Une preuve séquentielle d'indice n se traduit en un réseau d'indice n .

DÉMONSTRATION: Par induction sur la construction de la preuve. ◇

¶B.1. Elimination des Coupures

Donnons ici les étapes élémentaires:





PROPOSITION 5.9. *L'élimination des coupures est confluente et termine toujours.*

DÉMONSTRATION: Il suffit de remarquer que la taille du réseau (son nombre de liens) diminue au cours de chacune de ces étapes et que ce processus est localement confluente. $\diamond$

PROPOSITION 5.10. *L'élimination des coupures préserve l'indice.*

DÉMONSTRATION: Il suffit de vérifier que chaque étape élémentaire préserve l'indice, ce qui est évident (mais important). $\diamond$

§C. Séquentialisation

Les résultats du chapitre précédent s'étendent aisément à ce calcul-ci. Il faut tout de même vérifier qu'un réseau d'indice n admet une séquentialisation d'indice n . Pour ce faire, nous allons ici établir que chacune des étapes élémentaire de la séquentialisation décrites au chapitre précédent, tant par la méthode du **tenseur** scindant que

celle du **par** scindant préservent cet indice.

PROPOSITION 5.11. *Soit Π un réseau réduit à un lien axiome, hypothèse, ou $\perp$. Alors son indice est l'indice de sa séquentialisation.*

DÉMONSTRATION: Trivial. $\diamond$

PROPOSITION 5.12. *Soit Π un réseau connexe dont toutes les conclusions soient des tenseurs, t un lien tenseur scindant, Π_1 et Π_2 les deux parties du réseau $\Pi - t$ (Π_1 et Π_2 sont des réseaux car le lien tenseur t est scindant); soit π_1 et π_2 les deux preuves séquentielles correspondant à Π_1 et Π_2 et $\pi = \text{TENSEUR}(\pi_1, \pi_2)$ celle correspondant à Π ; si $\#(\pi_1) = \#(\Pi_1)$ et $\#(\pi_2) = \#(\Pi_2)$ alors $\#(\pi) = \#(\Pi)$.*

DÉMONSTRATION: Cela est dû aux égalités triviales suivantes:

$$\star \#(\pi) = \#(\pi_1) + \#(\pi_2)$$

$$\star cc(\Pi) = cc(\Pi_1) + cc(\Pi_2) - 1$$

$$\star ff(\Pi) = ff(\Pi_1) + ff(\Pi_2)$$

$\diamond$

PROPOSITION 5.13. *Soit Π un réseau non-connexe et soit Π_1 et Π_2 les classes d'une partition en deux classes des composantes connexes de Π (Π_1 et Π_2 sont évidemment des réseaux); soit π_1 et π_2*

les deux preuves séquentielles correspondant à Π_1 et Π_2 et $\pi = \text{MÉLANGE}(\pi_1, \pi_2)$ celle correspondant à Π ; si $\#(\pi_1) = \#(\Pi_1)$ et $\#(\pi_2) = \#(\Pi_2)$ alors $\#(\pi) = \#(\Pi)$.

DÉMONSTRATION: Cela est dû aux égalités triviales suivantes:

- * $\#(\pi) = \#(\pi_1) + \#(\pi_2) - 1$
- * $cc(\Pi) = cc(\Pi_1) + cc(\Pi_2)$
- * $ff(\Pi) = ff(\Pi_1) + ff(\Pi_2)$

◇

PROPOSITION 5.14. Soit Π un réseau, soit p l'un de ses liens par finaux; soit Π_1 le réseau obtenu par suppression de ce lien par final; soit π_1 la preuve séquentielle correspondant à Π_1 $\pi = \text{PAR}(\pi_1)$ celle correspondant à Π ; si $\#(\pi_1) = \#(\Pi_1)$ alors $\#(\pi) = \#(\Pi)$.

DÉMONSTRATION: Cela est dû aux égalités triviales suivantes:

- * $\#(\pi) = \#(\pi_1)$
- * $cc(\Pi) = cc(\Pi_1)$
- * $ff(\Pi) = ff(\Pi_1)$

◇

PROPOSITION 5.15. Soit Π un réseau connexe d'hypothèses $H_1, \dots, H_p$ et de conclusions $C_1, \dots, C_n$; soient p un lien par scindant, Π_1 et Π_2 les deux parties du réseau $\Pi - p$ (Π_1 et Π_2 sont des réseaux car le lien par t est scindant); on suppose que Π_1 d'indice i est d'hypothèses $H_1, \dots, H_i$ et de conclusions $C_1, \dots, C_k$ tandis que Π_2 d'indice j est d'hypothèses $H_{k+1}, \dots, H_p$ et de conclusions $C_{k+1}, \dots, C_n$ soit π_1 et π_2 les deux preuves séquentielles correspondant à Π_1 et Π_2 ; on obtient une séquentialisation π de Π ainsi:

- (i) on ajoute à tous les indice en dessous de l'hypothèse $A \wp B$ l'entier i puis on remplace dans π_2 la formule $A \wp B$ par la séquence $A \wp B, C_1, \dots, C_k$, et on appelle π'_2 le résultat de cette opération — qui n'est pas tout à fait une preuve séquentielle.
- (ii) on place la preuve séquentielle π_1 au dessus de la feuille $\vdash^i A \wp B, C_1, \dots, C_k$ — on obtient ainsi une preuve séquentielle d'indice $i + j$.

On a alors $\#(\Pi) = \#(\pi)$.

DÉMONSTRATION:

- * $cc(\Pi) = cc(\Pi_1) + cc(\Pi_2) - 1$
- * $ff(\Pi) = ff(\Pi_1) + ff(\Pi_2)$

◇

Toutes les étapes de la séquentialisation, que ce soit par la méthode du tenseur scindant (qui n'utilise pas la dernière proposition) ou par celle du par scindant préservent l'indice. D'où le

THÉORÈME 5.16. Soit Π un réseau d'indice n ; alors il existe une preuve séquentielle d'indice n tel que son réseau associé soit Π .

§D. η -Expansion

Ce calcul à la propriété de η -expansion: tout axiome non-atomique est dérivable à partir des seuls axiomes atomiques, et la démonstration en est aussi triviale qu'habituellement.

On remarquera que les deux preuves MÉLANGE(FAUX, UN) et AXIOME $\{\perp\}$ ont le même indice. En termes de réseaux, cela signifie que, *si cela ne crée pas de cycle bigarré*, on peut lier par un lien axiome n'importe quel lien $\perp$ à n'importe quel lien 1 sans changer l'indice du réseau.

Partie C

Réseaux et séquents ordonnés

Chapitre 6

Des espaces cohérents au connecteur "précède"

La sémantique cohérente est une sémantique de Heyting (on interprète les preuves et non les formules) étroitement liée à la logique linéaire dont l'étude a parfois suggéré de nouvelles directions. Ici, elle nous conduira à envisager un connecteur multiplicatif binaire, non-commutatif, associatif et autodual, **précède**, noté $\prec$, et un calcul logique dont les conclusions sont des multi-ensembles ordonnés de formules. Nous le développerons dans les chapitres suivants en gardant à l'esprit cette sémantique qui assure notamment le bien-fondé de notre calcul ordonné.

§A. Espaces cohérents

DÉFINITION 6.1. *Un espace cohérent A est un graphe simple non orienté dont l'ensemble des sommets ou points s'appelle la trame $|A|$. Deux sommets x et y sont dits strictement cohérents s'ils sont adjacents (et donc distincts); on note cela $x \frown y[A]$.*

Le dual d'un espace cohérent A , noté $A^\perp$ est le graphe simple complémentaire: A et $A^\perp$ ont la même trame et deux points distincts sont adjacents dans $A^\perp$ si et seulement s'ils ne le sont pas dans A :

$$\star |A^\perp| = |A|$$

$$\star \forall x, y \in |A^\perp| = |A| \quad x \neq y \Rightarrow (x \frown y[A^\perp] \iff \neg(x \frown y)[A])$$

On utilise les abréviations désormais standard de [GLT88, Gir87a]:

$x \frown y$	ssi	$x \neq y$ et x et y sont adjacents	x et y sont strictement cohérents
$x \times y$	ssi	$\neg(x \frown y)$	x et y sont incohérents
$x \smile y$	ssi	$\neg(x \frown y) \wedge x \neq y$	x et y sont strictement incohérents
$x \odot y$	ssi	$(x \frown y) \vee x = y$	x et y sont cohérents

Remarquons que la donnée de l'une de ces relations détermine totalement l'espace cohérent.

Les seuls morphismes entre espaces cohérents que nous envisagerons ici sont les application linéaires.

DÉFINITION 6.2. *Soit A un espace cohérent, on note $Cl(A)$ l'ensemble des cliques de A .*

Une application linéaire d'un espace cohérent A dans un espace cohérent B est une application l de $|A|$ dans $Cl(B)$ telle que:

$$\star \text{ si } x \odot y \text{ alors } l(x) \cup l(y) \in Cl(B)$$

$$\star \text{ si } x \wedge y \text{ alors } l(x) \cap l(y) = \emptyset$$

L'interprétation du calcul habituel procède ainsi:

On suppose que l'on dispose d'un espace cohérent pour chaque variable propositionnelle, et on définit alors, pour chaque formule, un espace cohérent lui correspondant.

On interprète alors chaque preuve du multi-ensemble $A_1, \dots, A_n$ comme une clique de l'espace cohérent $A_1 \wp \dots \wp A_n$.

Le calcul de la clique associée à une preuve s'effectue soit sur les réseaux (par la méthode des expériences que nous étudierons dans le cas plus général des réseaux ordonnés) soit inductivement sur la preuve séquentielle (nous donnerons ici la sémantique cohérente du calcul de la première partie, et celle du calcul ordonné lorsqu'il sera défini).

La sémantique d'une preuve est alors un invariant de l'élimination des coupures.

§B. Le connecteur précède

Soient A et B deux formules et $\heartsuit$ un connecteur binaire. Ce connecteur s'interprète dans la sémantique cohérente comme une fonction $\heartsuit$ qui associe au couple d'espace cohérents A, B un espace cohérent $A \heartsuit B$ dont la cohérence n'est fonction que de celle sur $|A|$ et de celle sur $|B|$.

Le connecteur $\heartsuit$ est dit multiplicatif si sa trame $|A \heartsuit B|$ est le produit cartésien $|A| \times |B|$ des trames de A et de B . Le connecteur $\heartsuit$ est donc en fait une loi de composition sur l'ensemble à trois éléments $\{ "\cup", "=", "\wedge" \}$.

Pour pouvoir associer aux cliques de A et de B des cliques de $A \heartsuit B$ de manière canonique il faut de plus que:

$$\star (x, y) \wedge (x', y') [A \heartsuit B] \iff x \wedge x' [A]$$

$$\star (x, y) \wedge (x, y') [A \heartsuit B] \iff y \wedge y' [B]$$

$$\star x \wedge x' [A] \wedge y \wedge y' [B] \Rightarrow (x, y) \wedge (x', y') [A \heartsuit B]$$

$$\star x \cup x' [A] \wedge y \cup y' [B] \Rightarrow (x, y) \cup (x', y') [A \heartsuit B]$$

Tout ceci peut se résumer ainsi: il faut que la loi de composition sur $\{ "\cup", "=", "\wedge" \}$ soit compatible avec l'ordre $"\cup" < "=" < "\wedge"$ dans le sens suivant: $s \leq s' \wedge t \leq t' \Rightarrow s \heartsuit s' \leq t \heartsuit t'$. Compte tenu de l'implication $(s \heartsuit s' = "=") \Rightarrow (s = "=" \wedge s' = "=")$ on voit que seuls $"\wedge \heartsuit \cup"$ et $"\cup \heartsuit \wedge"$ ne sont pas déterminés par cette condition.

On dira que le connecteur $\heartsuit$ est commutatif ssi

$$(x, y) \odot (x', y') [A \heartsuit B] \iff (y, x) \odot (y', x') [B \heartsuit A]$$

i.e. si la loi de composition correspondante l'est. On voit alors qu'il n'y a que deux connecteurs multiplicatifs binaires commutatifs, que nous appellerons bien évidemment **par** et **tenseur** qui correspondent aux formules du même nom:

$A \heartsuit B$			
$A \backslash B$	$\smile$	$=$	$\frown$
$\smile$	$\smile$	$\smile$	$\frown$
$=$	$\smile$	$=$	$\frown$
$\frown$	$\frown$	$\frown$	$\frown$

$A \otimes B$			
$A \backslash B$	$\smile$	$=$	$\frown$
$\smile$	$\smile$	$\smile$	$\smile$
$=$	$\smile$	$=$	$\frown$
$\frown$	$\smile$	$\frown$	$\frown$

soit

$A \otimes B$	$(a, b) \odot (a', b')$	$[A \otimes B]$	ssi	$a \odot a'[A] \wedge b \odot b'[B]$
$A \heartsuit B$	$(a, b) \frown (a', b')$	$[A \heartsuit B]$	ssi	$a \frown a'[A] \vee b \frown b'[B]$

et on a bien $(A \heartsuit B)^\perp = A^\perp \otimes B^\perp$.

S'il on se demande quels sont les connecteurs multiplicatifs binaires non-commutatifs on voit qu'il n'y en a que deux, qui sont symétriques l'un de l'autre. On appellera ces connecteurs A "précède" B et B "précède" A ce que l'on notera $A < B$ et $B < A$. En voici la définition:

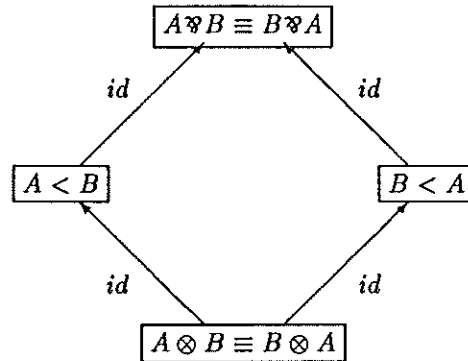
$A < B$			
$A \backslash B$	$\smile$	$=$	$\frown$
$\smile$	$\smile$	$\smile$	$\frown$
$=$	$\smile$	$=$	$\frown$
$\frown$	$\smile$	$\frown$	$\frown$

Ce qui s'écrit aussi:

$A < B$	$(a, b) \frown (a', b')$	$[A < B]$	ssi	$(a \frown a'[A] \wedge b = b') \vee b \frown b'[B]$
---------	--------------------------	-----------	-----	--

PROPOSITION 6.3. *Le connecteur précède est*

- (i) *non-commutatif* $A < B \not\equiv B < A$
- (ii) *associatif i.e.* $A < (B < C) \equiv (A < B)C$
- (iii) *autodual i.e.* $(A < B)^\perp \equiv A^\perp < B^\perp$
- (iv) *"entre le par et le tenseur": on a $(A \otimes B) \multimap (A < B)$ et, par dualité, $(A < B) \multimap (A \heartsuit B)$, soit:*



DÉMONSTRATION:

(i) Cela se voit sur la définition de la cohérence.

(ii) Montrons que:

- (1) $(x, (y, z)) \wedge (x', (y', z')) [A < (B < C)]$ équivaut à
 (2) $((x, y), z) \wedge ((x', y'), z') [(A < B) < C]$, ce qui suffira puisque $(x, (y, z)) = (x', (y', z'))$ ssi $((x, y), z) = ((x', y'), z')$. Les lignes suivantes sont équivalentes:

$$\begin{aligned}
 (1) : & (x, (y, z)) \wedge (x', (y', z')) [A < (B < C)] \\
 & (x \wedge x' [A] \wedge (y, z) = (y', z')) \vee ((y, z) \wedge (y', z') [B < C]) \quad (def) \\
 & (x \wedge x' [A] \wedge y = y' \wedge z = z') \vee ((y \wedge y' [B] \wedge z = z') \vee z \wedge z' [C]) \quad (def) \\
 & (x \wedge x' [A] \wedge y = y' \wedge z = z') \vee (y \wedge y' [B] \wedge z = z') \vee z \wedge z' [C] \quad (ass) \\
 & (x \wedge x' [A] \wedge y = y' \wedge z = z') \vee (y \wedge y' [B] \wedge z = z') \vee z \wedge z' [C] \quad (dist) \\
 & \left(\left((x \wedge x' [A] \wedge y = y') \vee y \wedge y' \right) \wedge z = z' \right) \vee z \wedge z' \quad (def) \\
 & ((x, y) \wedge (x', y') [A < B] \wedge z = z') \vee z \wedge z' \quad (def)
 \end{aligned}$$

$$(2) : ((x, y), z) \wedge ((x', y'), z') [(A < B) < C]$$

(iii) $(x, y) \wedge (x', y') [A < B]$ signifie, d'après la table, $y \wedge y' [B] \vee (x \wedge x' [A] \wedge y \wedge y' [B])$ soit $y \wedge y' [B^\perp] \vee (x \wedge x' [A^\perp] \wedge y \wedge y' [B^\perp])$ i.e. $(x, y) \wedge (x', y') [A^\perp < B^\perp]$

(iv) L'identité de $A \otimes B$ dans $A < B$ une application linéaire: les tables de cohérence précédentes montrent que $(x, y) \odot (x', y') [A \otimes B] \Rightarrow (x, y) \odot (x', y') [A < B]$, ce qui suffit.

◇

§C. Produit ordonné d'espaces cohérents

La définition de la cohérence dans l'espace cohérent $A < B$ donne l'idée de la généralisation suivante:

DÉFINITION 6.4. (produit ordonné d'espaces cohérents) Soit $\mathbf{i}$ un ordre sur le multi-ensemble d'espaces cohérents $A_1, \dots, A_n$. On définit le produit ordonné par $\mathbf{i}$ des espaces cohérents $A_1, \dots, A_n$

ainsi:

★ Sa trame est le produit cartésien des espaces cohérents:

$$\left| \prod_i A_i \right| = \prod_{i \in I} |A_i|$$

★ La cohérence est définie "lexicographiquement":

$$(a_1, \dots, a_n) \frown (a'_1, \dots, a'_n) \left[\prod_i A_i \right] \iff \exists i \left[a_i \frown a'_i[A_i] \wedge \left(\forall j > i \ [i] \ a_j = a'_j[A_j] \right) \right]$$

On remarque alors que $A \not\approx B$ n'est autre que $\prod_{\emptyset} A, B$ tandis que $A < B$ n'est autre que $\prod_{A < B} A, B$.

Cette définition suggère qu'on peut manipuler des multi-ensembles ordonnés de formules au lieu des multi-ensembles habituels, i.e. travailler avec des connecteurs multiplicatifs généralisés (qui ne seront pas tous définissables avec le *par* et le *précède*). C'est là l'objet des chapitres suivants.

§D. Sémantique cohérente du calcul multiplicatif avec mélange

A titre d'exemple, nous donnons rapidement la sémantique cohérente du calcul des séquents défini dans le chapitre 4.

π	$\ \pi\ $
$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \\ \vdots \end{array} \quad \vdash A_1, A_2, \dots, A_n \quad \vdash B_1, B_2, \dots, B_p}{\vdash A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p} \text{MÉLANGE}$	$\left\{ (a_1, \dots, a_n, b_1, \dots, b_p) / \right. \\ \left. (a_1, \dots, a_n) \in \ \pi_1\ \text{ et } (b_1, \dots, b_p) \in \ \pi_2\ \right\}$
$\overline{\vdash A, A^\perp} \text{ AXIOME}$	$\left\{ (a, a) / a \in A \right\}$
$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \\ \vdots \end{array} \quad \vdash A_1, A_2, \dots, A_n, K \quad \vdash K^\perp, B_1, B_2, \dots, B_p}{\vdash A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p} \text{COUPURE}$	$\left\{ (a_1, \dots, a_n, b_1, \dots, b_p) / \exists z \in K \right. \\ \left. (a_1, \dots, a_n, z) \in \ \pi_1\ \text{ et } (b_1, \dots, b_p, z) \in \ \pi_2\ \right\}$
$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \\ \vdots \end{array} \quad \vdash A_1, A_2, \dots, A_n, A \quad \vdash B, B_1, B_2, \dots, B_p}{\vdash A_1, A_2, \dots, A_n, A \otimes B, B_1, B_2, \dots, B_p} \text{TENSEUR}$	$\left\{ (a_1, \dots, a_n, (a, b), b_1, \dots, b_p) / \right. \\ \left. (a_1, \dots, a_n, a) \in \ \pi_1\ \text{ et } (b, b_1, \dots, b_p) \in \ \pi_2\ \right\}$
$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \end{array} \quad \vdash A, B, A_1, A_2, \dots, A_n}{\vdash A \wp B, A_1, A_2, \dots, A_n} \text{PAR}$	$\left\{ (a_1, \dots, a_n, (a, b)) / \right. \\ \left. (a_1, \dots, a_n, a, b) \in \ \pi_1\ \right\}$

On vérifie alors aisément que la sémantique cohérente d'une preuve est un invariant de la preuve par rapport à l'élimination des coupures.

On pourrait tout aussi bien calculer la sémantique cohérente d'une preuve sur le réseau correspondant. Ce calcul étant une restriction du calcul ordonné, ce sera un cas particulier de la sémantique cohérente des réseaux ordonnés présentée ultérieurement.

Chapitre 7

Réseaux ordonnés

Nous allons donc construire un système de preuves dont la conclusion est un multi-ensemble ordonné de formules construites à l'aide de **par**, **tenseur** et **précède**. La sémantique de Heyting d'une preuve de $A_1, \dots, A_n$ dans l'ordre i sera une clique de l'espace cohérent produit précédemment défini. On commence par présenter ces preuves en réseaux car la généralisation correspondante des réseaux est plus naturelle que celle du calcul des séquents.

Les formules sont ici écrites à l'aide des trois connecteurs binaire **tenseur** " $\otimes$ ", **par** " $\wp$ ", et **précède** " $<$ ". On rappelle que $(A < B)^\perp = A^\perp < B^\perp$.

§A. Définition des réseaux ordonnés

NOTATION 7.1. [couleurs] On se donne un ensemble dénombrable de couleurs. Parmi ces couleurs, on en distingue une, le noir, notée $\mathcal{N}$ tandis que les autres seront dites propres, et seront désignées par des lettres grecques $\alpha, \beta, \dots$

NOTATION 7.2. Par graphe simple orienté on entend un graphe orienté sans boucle tel que, pour tout couple de points (x, y) on ait au plus un arc de x vers y , ce qui n'interdit pas d'avoir aussi un arc de y vers x . Si la paire d'arcs $u = (x, y)$ et $\tilde{u} = (y, x)$ font partie du graphe, on parlera alors de l'arête $\{x, y\}$ ce qu'on écrira souvent $x \text{---} y$. En fait on réservera l'expression l'arc (x, y) , ce qu'on écrira souvent $x \rightarrow y$ pour un arc $u = (x, y)$ du graphe tel que l'arc $\tilde{u} = (y, x)$ ne fasse pas partie du graphe. Dire qu'une arête est de couleur γ , c'est dire que les deux arcs qui la constituent sont de couleur γ .

On va définir les réseaux ordonnés comme des graphes simples orientés colorés, dont les arcs sont tous noirs et dont les arêtes sont colorées (et dont certaines sont noires).

DÉFINITION 7.3. On appelle (très abusivement) arbre des sous-formules d'une formule F de ce langage, un couple (T, A) où

- * T est un arbre binaire dont les sommets sont étiquetés par des sous-formules de F , et les branchements par les connecteurs (binaires), " $\wp$ ", " $\otimes$ " et " $<$ ", de sorte que: si G et H sont les successeurs de L dans un branchement étiqueté par le connecteur $\wp$ alors $L = G \wp H$. On

remarquera qu'on s'autorise ainsi à considérer des arbres des sous formules "élagués" dont les feuilles ne sont pas toujours des atomes; par contre, si une sous-formule immédiate H , d'une formule $H \heartsuit G$ figure dans l'arbre, l'autre sous-formule immédiate G y figure aussi.

- * A est l'ensemble des arcs $G \rightarrow H$ pour tous les couples $G H$ de sommets de l'arbre dont le prédécesseur est $G < H$.

Ce n'est donc pas un arbre, mais un graphe simple constitué par un arbre, augmenté d'arcs entre certains sommets.

DÉFINITION 7.4. Une famille d'arbres des sous-formules est dite bien colorée si et seulement si:

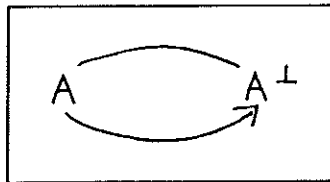
- * les arcs sont noirs
- * les deux arêtes correspondant à un connecteur tenseur sont noires;
- * les deux arêtes correspondant à un même connecteur par ou précède sont d'une même couleur propre;
- * deux arêtes appartenant à des connecteurs par ou précède différents, — qu'ils soient ou non dans le même arbre — sont de couleurs (propres) différentes.

DÉFINITION 7.5. Un préréseau ordonné est un graphe simple coloré dont les sommets sont étiquetés par des formules et le signe " $\bullet$ ". Plus précisément, un préréseau ordonné Π est défini par un quadruplet $(\mathcal{F}, A\mathbf{x}, \text{Cut}, \mathbf{i})$ où:

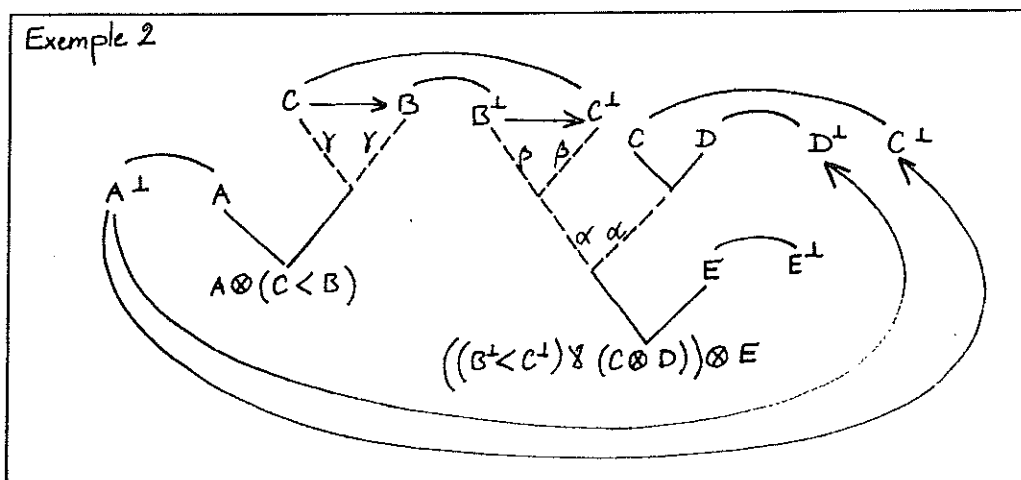
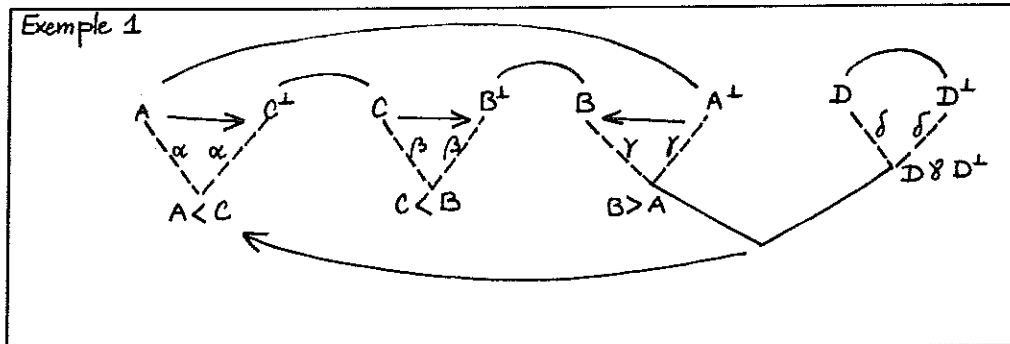
- * $\mathcal{F}$ est une famille bien colorée d'arbres des sous-formules des formules F_i , dont les feuilles sont les A_i .
- * $A\mathbf{x}$ est un ensemble d'arêtes noires non-adjacentes de la forme $A_i A_j$ où $A_i = A_j^\perp$
- * Cut est un ensemble triplets $(\bullet_{\mathbf{k}}, c_{\mathbf{k}}, c_{\mathbf{k}}^\perp)$ où:
 - ▷ les $\bullet_{\mathbf{k}}$ sont des sommets du préréseau ordonné autres que ceux de $\mathcal{F}$, étiquetés $\bullet$, de degré 2, appelés coupures,
 - ▷ pour tout $\mathbf{k}$, $c_{\mathbf{k}}$ et $c_{\mathbf{k}}^\perp$ sont des arêtes noires de la forme $\bullet_{\mathbf{k}} \text{---} F_i$ et $\bullet \text{---} F_j$ avec $F_i = F_j^\perp$, de sorte que, si $\mathbf{k} \neq \mathbf{k}'$, les deux arêtes $c_{\mathbf{k}}$ et $c_{\mathbf{k}'}$ ne sont pas adjacentes et les deux arêtes $c_{\mathbf{k}}$ et $c_{\mathbf{k}'}$ ne sont pas adjacentes. (On peut aussi dire qu'une racine F_i est incidente à au plus une arête $c_{\mathbf{k}}$ ou $c_{\mathbf{k}}^\perp$)
- * $\mathbf{i}$ est un ensemble d'arcs noirs dont les extrémités sont parmi $F_i, \bullet_{\mathbf{k}}$

On appelle **feuille** du préréseau ordonné toute feuille de la pseudo-forêt $\mathcal{F}$. On appelle **coupures** du préréseau ordonné les sommets de Cut (i.e. les $\bullet_{\mathbf{k}}$). On appelle **conclusion** du préréseau ordonné toute racine F_i de $\mathcal{F}$ non-incidente à un sommet de Cut . On appelle **hypothèse** d'un préréseau ordonné une feuille non-incidente à une arête de $A\mathbf{x}$.

Remarquons tout de suite que l'objet suivant n'est même pas un préréseau ordonné (ce n'est pas un graphe simple orienté):

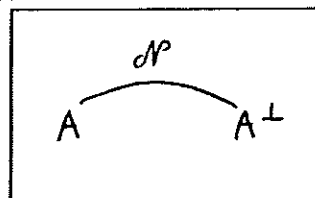


et donnons maintenant deux exemples de préreseaux ordonnés:

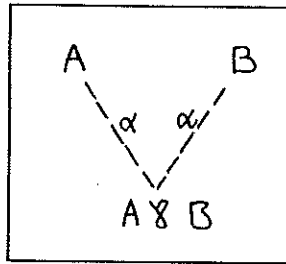


DÉFINITION 7.6. Un lien d'un préreseau ordonné $\Pi = (\mathcal{F}, Ax, Cut, i)$ est l'un de ses (petits) sous-graphes pleins suivants:

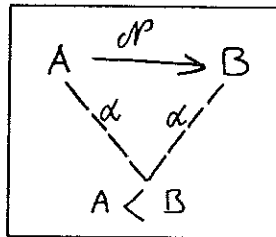
- * Un lien hypothèse est le graphe réduit à une hypothèse du préreseau ordonné.
- * Un lien axiome est une arête de Ax .



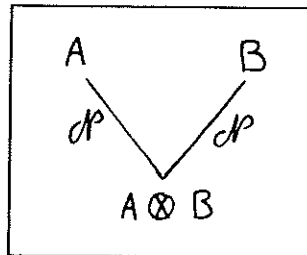
- * Un lien par est constitué des deux arêtes incidentes d'une même couleur propre liant une sous-formule $A \wp B$ de F_i à ses deux sous-formules immédiates A et B . $A \wp B$ est appelée la conclusion du lien, tandis que ses sous formules immédiates A et B sont appelées prémisses de ce lien.



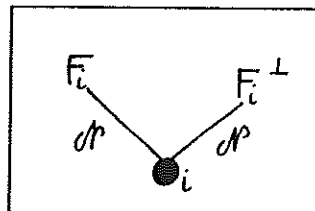
- * Un **lien précède** est constitué des deux arêtes incidentes d'une même couleur propre liant une sous-formule $A < B$ de F_i à ses deux sous-formules immédiates A et B , et de l'arc noir $A \rightarrow B$. $A < B$ est appelée la conclusion du lien, tandis que ses sous formules immédiates A et B sont appelées prémisses de ce lien.



- * Un **lien tenseur** est constitué des deux arêtes noires liant une sous-formule $A \otimes B$ de F_i à ses deux sous-formules immédiates A et B . $A \otimes B$ est appelée la conclusion du lien, tandis que ses sous formules immédiates A et B sont appelées prémisses de ce lien.



- * Un **lien coupure** est constitué des deux arêtes noires incidentes à un sommet de Cut.



Un préréseau ordonné s'obtient donc, s'il n'y a pas de coupure ¹, à partir d'un préréseau commutatif par addition d'arcs noirs de deux sortes. Les premiers vont d'une prémisses d'un lien par à l'autre, et les second vont d'une conclusion à une autre conclusion.

DÉFINITION 7.7. Dans un graphe coloré orienté on dit qu'un ensemble d'arcs et d'arêtes est **bigarré** s'il ne contient pas deux arêtes d'une même couleur propre. Un chemin bigarré est donc

¹sinon il faudrait que réseau commutatif contienne des coupures entre deux $\otimes$; on verra au paragraphe §F. de ce chapitre le rapport exact entre réseaux ordonnés et réseaux commutatifs

un chemin qui n'emprunte jamais deux arêtes d'une même couleur propre — ce qui ressemble, si ce n'était de l'orientation, aux définitions similaires des chapitre 2 et 4.

DÉFINITION 7.8. On dit qu'un *préréseau ordonné* est un *réseau ordonné* s'il ne contient pas de *circuit bigarré*.

REMARQUE 7.9. Dans un *réseau ordonné*, $\Pi = (\mathcal{F}, Ax, Cut, i)$ le sous-graphe noir i est sans circuit. Un *préréseau* $\Pi = (\mathcal{F}, Ax, Cut, i)$ est un *réseau ordonné* si et seulement si le *préréseau* $\Pi' = (\mathcal{F}, Ax, Cut, \bar{i})$ est un *réseau ordonné* où $\bar{i}$ désigne la clôture transitive de i (qui est antiréflexive puisque i est antiréflexive et sans circuit).

Ainsi le second exemple de *préréseau* est-il un *réseau*, tandis que le premier n'en est pas un.

DÉFINITION 7.10. Soient $\Pi = (\mathcal{F}, Ax, Cut, i)$ un *préréseau de conclusions* $A_1, \dots, A_n, \bullet_1, \dots, \bullet_k[i]$. On note Π^- le *préréseau* $\Pi^- = (\mathcal{F}, Ax, Cut, \emptyset)$; il est clair que si Π est un *réseau* alors Π^- est a fortiori un *réseau*.

On appelle *graphe de correction* de Π tout sous-graphe bigarré maximal de Π . Les graphes de corrections des *réseaux ordonnés* sont donc orientés.

Un sous-graphe maximal de Π^- qui est bigarré sera dit *graphe de correction interne*.

On appelle *parcours* de Π toute relation u sur $A_1, \dots, A_n$ définie par un *graphe de correction interne*, i.e. u est un *parcours* de R si et seulement si il existe un *graphe de correction interne* G tel que:

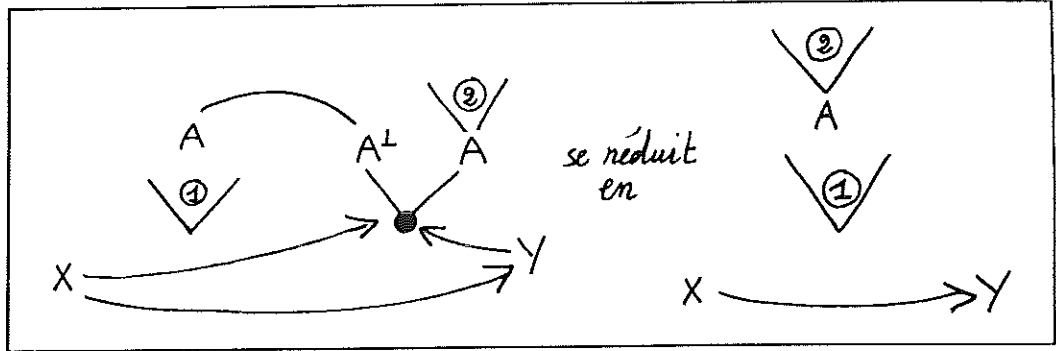
$A_i u A_k$ si et seulement si G contient un chemin de A_i à A_k . Les *parcours* ne sont donc pas forcément symétriques, contrairement aux *parcours* habituels. Ils sont par contre transitifs.

PROPOSITION 7.11. Un *préréseau ordonné* de conclusions $A_1, \dots, A_n[u]$ est un *réseau ordonné* si et seulement si:

- * i est sans circuit
- * tout *graphe de correction interne* est sans circuit
- * i est orthogonal à tout *parcours interne* (rappelons que l'orthogonalité est stable par clôture transitive et réflexive)

§B. Elimination des Coupures Ordonnées

Nous allons maintenant démontrer que le critère donné dans la définition des *réseaux ordonnés* est stable par élimination des coupures. Nous donnons ici les étapes élémentaires de l'élimination des coupures, qui utilisent les expansions d'ordres présentées au chapitre 1, et nous montrons que chacune de ces étapes donne un *réseau ordonné*, i.e. que le critère *sans circuit bigarré* est préservé par élimination des coupures. Notons que c'est là un argument essentiel pour pouvoir donner le nom de calcul à ce système, et que faute de ce résultat nous aurions soit modifié ce calcul, soit abandonné notre idée d'un tel calcul.

¶B.1. Coupure sur un axiome $[ax]$ 

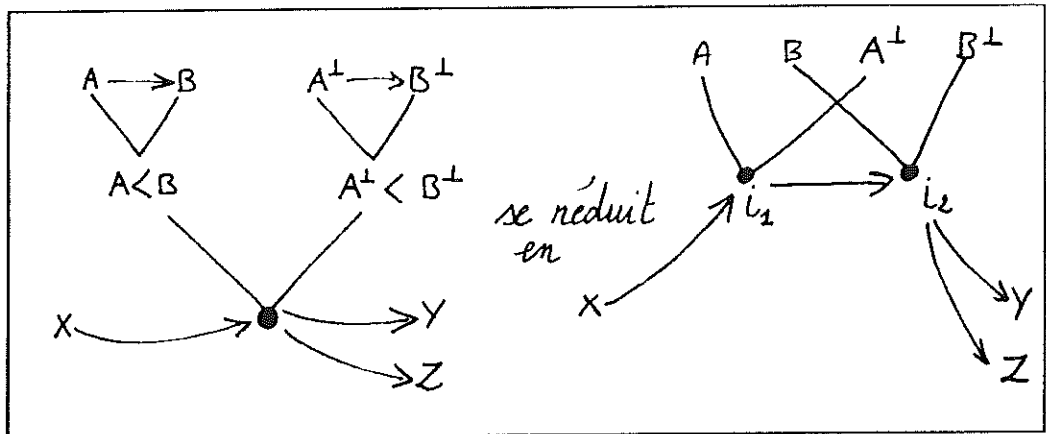
L'ordre sur les conclusions et coupures $\Gamma, \bullet$ est $i|_{\Gamma}$.

Montrons maintenant que:

LEMME 7.12. Si Π est un réseau ordonné qui se réduit en une étape $[ax]$ en Π' , alors Π' est aussi un réseau ordonné.

DÉMONSTRATION: Il est clair que Π' est un pré-réseau, et qu'il ne contient pas de circuit bigarré interne.

Il suffit alors de remarquer que Π' s'obtient à partir de Π par contraction d'arêtes noires (à double sens), et par effacement d'arcs noirs de i . $\diamond$

¶B.2. Coupure entre deux précédé $[pcd]$ 

L'ordre sur les conclusions et coupures est $i' = i /_{\bullet_i \rightsquigarrow (\bullet_{i_1}, \prec \bullet_{i_2})}$

LEMME 7.13. Si Π est un réseau ordonné qui se réduit en une étape $[pcd]$ en Π' , alors Π' est aussi un réseau ordonné.

DÉMONSTRATION: Il est clair que Π' est un pré-réseau et que tout chemin bigarré de Π' utilisant au plus une arête noire de chaque paires d'arêtes $\bullet_{i_1} \text{---} A, \bullet_{i_2} \text{---} B$ et $\bullet_{i_1} \text{---} A^{\perp}, \bullet_{i_2} \text{---} B^{\perp}$ se relève en un chemin bigarré de Π .

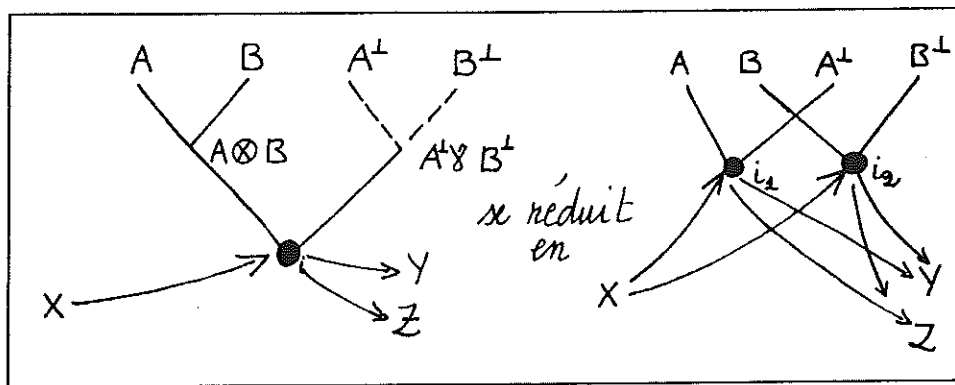
Supposons donc que Π' contienne un circuit bigarré; il en contient alors un qui est de longueur minimale, ce qu'on abrègera en minimal.

D'après la remarque précédente, si ce circuit bigarré minimal c contient au plus une arête de chaque paires d'arêtes noires $\bullet_{i_1} \text{---} A, \bullet_{i_2} \text{---} B$ et $\bullet_{i_1} \text{---} A^\perp, \bullet_{i_2} \text{---} B^\perp$, ce circuit se relève en un circuit bigarré de Π .

Supposons alors que c emprunte la paire d'arêtes $\bullet_{i_1} \text{---} A, \bullet_{i_2} \text{---} B$. Puisque c est minimal, c emprunte nécessairement l'arc $\bullet_{i_1} \rightarrow \bullet_{i_2}$. Il existe donc dans Π' — et donc dans Π — un chemin bigarré de B à A . Ce chemin et l'arc $A \rightarrow B$ constituent un chemin bigarré de Π .

Le cas où c contient la paire d'arêtes $\bullet_{i_1} \text{---} A^\perp, \bullet_{i_2} \text{---} B^\perp$ se traite de manière similaire. $\diamond$

¶B.3. Coupure entre un tenseur et un par $[ts/par]$



LEMME 7.14. Si Π est un réseau ordonné qui se réduit en une étape $[ts/par]$ en Π' , alors Π' est aussi un réseau ordonné.

L'ordre sur les conclusions et coupures est $i' = i / \bullet_{i_1} \sim \bullet_{i_2}$.

DÉMONSTRATION: Il est clair que Π' est un pré-réseau et que tout chemin bigarré de Π' utilisant au plus une arête noire de la paire d'arêtes $\bullet_{i_1} \text{---} A^\perp, \bullet_{i_2} \text{---} B^\perp$ se relève en un chemin bigarré de Π .

Supposons donc que Π' contienne un circuit bigarré; il en contient alors un qui est de longueur minimale, ce qu'on abrègera en minimal. D'après la remarque précédente, si ce circuit bigarré minimal c contient au plus une arête de la paire d'arêtes noires $\bullet_{i_1} \text{---} A^\perp, \bullet_{i_2} \text{---} B^\perp$, ce circuit se relève en un chemin bigarré de Π . Supposons maintenant que ce circuit contienne la paire d'arêtes noires $\bullet_{i_1} \text{---} A^\perp, \bullet_{i_2} \text{---} B^\perp$.

S'il emprunte les deux arêtes $\bullet_{i_1} \text{---} A$ et $\bullet_{i_2} \text{---} B$, alors Π' — et donc Π — contient un chemin bigarré de A à B , qui constitue, avec les deux arêtes noires du lien tenseur $A \otimes B$ un circuit bigarré de Π .

S'il emprunte une seule des deux arêtes $\bullet_{i_1} \text{---} A, \bullet_{i_2} \text{---} B$, alors Π' — et donc Π — contient un chemin bigarré de l'une des formules $A^\perp, B^\perp$ à l'une des formules A, B (ou de l'une des formules A, B à l'une des formules $A^\perp, B^\perp$), qui, concaténé avec le chemin de Π passant par $\bullet_i$ les reliant, constitue un circuit bigarré de Π . $\diamond$

¶B.4. Normalisation forte et confluence du calcul des réseaux ordonnés

Nous allons montrer que l'élimination des coupures est une normalisation forte et confluente. Pour cela il faut préciser ce que l'on entend par élimination des coupures: soit Π un réseau ordonné de conclusion $A_1, \dots, A_n, \bullet_1, \dots, \bullet_p$ dans l'ordre i alors on peut, par les transformations locales $[ax]$, $[pcd]$ et $[ts/par]$ précédemment décrites, obtenir une preuve de $A_1, \dots, A_n$ dans l'ordre $i|_{A_1, \dots, A_n}$.

THÉORÈME 7.15. *Si Π est un réseau ordonné alors Π est fortement normalisable, et la normalisation est confluente.*

DÉMONSTRATION: On a vu que si Π est un réseau, alors tout réduit Π' de Π est aussi un réseau, i.e. ne contient pas de circuit bigarré (il est clair que Π' est un préréseau).

Montrons maintenant que que $i'|_{A_1, \dots, A_n} = i|_{A_1, \dots, A_n}$. Dans le premier cas $i' = i/\bullet_{i_1} \sim \bullet_{i_2}$, et dans le second, $i' = i/\bullet_{i_1} \sim \bullet_{i_2}$. D'après le paragraphe §C. du chapitre 1, la restriction de i' au conclusions et aux coupures autres que $\bullet_i$ est i , dans les deux cas, ce qui assure, a fortiori, le résultat.

Chacune des étapes élémentaires diminuant la taille du réseau, la terminaison est assurée. La confluence est triviale. $\diamond$

§C. Sémantique cohérente des réseaux ordonnés

Nous allons ici calculer la sémantique cohérente d'un réseau ordonné, par une méthode qui généralise à notre calcul les expériences de [Gir87a].

Tout d'abord une remarque aussi triviale qu'utile:

REMARQUE 7.16. Si $(a_1, \dots, a_n) \sim (a'_1, \dots, a'_n) \left[\prod_i A_i \right]$ et $a_k \sim a'_k [A_k]$ alors $\exists l > k[i] \quad a_l \sim a'_l [A_l]$.

(Evident d'après la définition du produit ordonné d'espaces cohérents.)

DÉFINITION 7.17. Une expérience est l'attribution à chacune des formules du réseau d'une valeur prise dans la trame de l'espace cohérent correspondant de sorte que:

- * pour chaque lien axiome ou coupure on ait la même valeur pour A et $A^\perp$
- * pour chaque lien par tenseur ou précède la valeur de la conclusion du lien soit le couple des valeurs des prémisses.

LEMME 7.18. (de compatibilité) Deux expériences d'un même réseau sont toujours cohérentes modulo le produit ordonné des conclusions.

DÉMONSTRATION: On suppose qu'on a réalisé deux expériences incohérentes, et on montre que le réseau n'est pas correct, i.e. contient un circuit bigarré. Si c'est le cas, il existe une conclusion où les deux expériences sont strictement incohérentes. En effet si les expériences

étaient toujours cohérentes ou égales sur les conclusions, comme elles ne sont pas égales en toutes les conclusions, en considérant une conclusion maximale où elles soient strictement cohérentes on voit qu'elles seraient cohérentes dans le produit ordonné des conclusions.

On construit, en partant de cette conclusion où elles sont strictement incohérentes, un chemin bigarré qui a la propriété suivante:

- * lorsque le chemin monte les expériences sont strictement incohérentes,
- * lorsque le chemin descend les expériences sont strictement cohérentes,
- * lorsque le chemin emprunte un axiome on passe d'une formule où les expériences étaient strictement incohérentes à une formule où elles sont strictement cohérentes,
- * lorsque le chemin emprunte une coupure on passe d'une formule où les expériences étaient strictement cohérentes à une formule où elles sont strictement incohérentes,
- * lorsque le chemin emprunte un arc de l'ordre, ou d'un lien précède, on passe d'une formule où les expériences étaient strictement cohérentes à une formule où elles sont strictement cohérentes.

On va montrer qu'on peut prolonger indéfiniment un tel chemin bigarré, à moins qu'il y ait un circuit bigarré. Cela assurera l'existence d'un circuit bigarré et contredira la correction du réseau. Pour un réseau correct, les valeurs des expériences en les conclusions sont donc nécessairement cohérentes en le produit ordonné des conclusions.

On envisage successivement (et patiemment) tous les lieux possibles de la fin du chemin déjà construit.

DANS UN LIEN par

EN MONTANT (les deux expériences sont donc strictement incohérentes en sa conclusion)

Un rapide calcul montre que les deux expériences sont strictement incohérentes en l'une au moins des deux prémisses. On étend le chemin jusqu'à cette prémisse, en s'assurant que le chemin n'a pas utilisé auparavant l'autre arête du lien par. Si tel était le cas, vu les cohérences, on y serait passé en montant, et on aurait donc un circuit bigarré.

EN DESCENDANT PAR L'UNE DES PRÉMISSES (les deux expériences sont donc strictement cohérentes en cette prémisse)

Les deux expériences sont aussi strictement cohérentes en la conclusion du lien. On prolonge notre chemin en descendant de la prémisse à la conclusion du lien, en remarquant que si notre chemin a déjà utilisé l'autre arête on a un circuit bigarré utilisant l'arête que l'on vient d'ajouter au chemin.

DANS UN LIEN précède : $A < B$

EN MONTANT (les deux expériences sont donc strictement incohérentes en sa conclusion)

Un rapide calcul montre que les deux expériences sont strictement incohérentes en l'une au moins des deux prémisses. Notons que, puisque les deux expériences sont strictement incohérentes en la conclusion, si le chemin a emprunté l'une des deux arêtes du lien, c'est en montant, et que de fait on a circuit bigarré. Si le chemin n'a

emprunté aucune de ces deux arêtes, on peut le prolonger jusqu'à une prémisse où les deux expériences sont strictement incohérentes.

EN DESCENDANT (les deux expériences sont donc strictement cohérentes en cette prémisse)

PAR LA PRÉMISSSE A Si on arrive à la prémisse A d'un lien précède $A < B$ en descendant les deux cas suivants peuvent se présenter:

* soit $B : \sim$ auquel cas on remonte par B en utilisant l'arc du lien précède ,

* soit $B : \circ$ auquel cas les expériences sont strictement cohérentes en la conclusion du lien et on prolonge le chemin en descendant jusqu'à la conclusion du lien. Si le chemin avait emprunté l'arête droite (forcément en descendant), alors il y a un circuit bigarré utilisant l'arc du lien précède .

PAR LA PRÉMISSSE B Dans ce cas les deux expériences sont strictement cohérentes en la conclusion de ce lien. On prolonge le chemin en descendant à cette conclusion. Notons que si le chemin avait emprunté l'arête $A \text{---} A < B$ du lien précède , c'était en descendant, et que dans ce cas on obtient un circuit bigarré empruntant l'arête droite du lien précède .

DANS UN LIEN tenseur

EN MONTANT (les deux expériences sont donc strictement incohérentes en sa conclusion)

Si on arrive en montant à la conclusion d'un lien tenseur , alors les deux expériences sont strictement incohérentes en l'une au moins des prémisses du lien. On prolonge notre chemin jusqu'à une telle prémisse.

EN DESCENDANT (les deux expériences sont donc strictement cohérentes en cette prémisse)

Si on arrive en descendant dans l'une des prémisses d'un lien tenseur , alors soit on remonte par l'autre prémisse si les deux expériences y sont strictement incohérentes, soit on descend à la conclusion où les deux expériences sont alors strictement cohérentes. (Un rapide calcul montre qu'il n'y a pas d'autre possibilité.)

EN MONTANT DANS UN LIEN axiome (les deux expériences sont donc strictement incohérentes en cette conclusion du lien axiome)

Si on arrive (nécessairement en montant) à l'une des conclusions d'un lien axiome on prolonge le chemin en descendant par l'autre conclusion du lien axiome (où les deux expériences sont forcément strictement cohérentes).

EN DESCENDANT DANS UN LIEN coupure (les deux expériences sont donc strictement cohérentes en cette prémisse du lien coupure)

Si on arrive (forcément en descendant) dans l'une des prémisses d'un lien coupure alors on prolonge notre chemin en remontant par l'autre prémisse du lien coupure (où les expériences sont forcément strictement incohérentes).

EN DESCENDANT DANS UNE CONCLUSION DU RÉSEAU (les deux expériences sont donc strictement cohérentes en cette conclusion)

Si on arrive (nécessairement en descendant) dans une conclusion du réseau comme nos deux expériences sont strictement incohérentes modulo le produit ordonné des conclusions, il existe une conclusion supérieure dans l'ordre sur les conclusions où les deux expériences sont strictement incohérentes (d'après la remarque 7.16.). Il existe donc un arc de notre conclusion à cette conclusion, arc par lequel on prolonge notre chemin.

Le lecteur méticuleux aura remarqué qu'on n'utilise pas l'ordre entre les coupures. Cela est

dû au fait que les coupures s'éliminent sans modifier la sémantique et que si cet ordre était nécessaire, on ne pourrait traduire notre chemin dans le réseau normalisé. $\diamond$

LEMME 7.19. *Tout réseau normal admet une sémantique cohérente non-triviale.*

DÉMONSTRATION: Il suffit de choisir, pour tout lien axiome $A, A^\perp$ une valeur dans la trame $|A| = |A^\perp|$ de l'espace cohérent correspondant à A . Ensuite, on "propage" ces valeurs dans le réseau, ce qui fournit toujours une expérience puisqu'il n'y a pas de coupure. $\diamond$

LEMME 7.20. *Si Π se réduit en Π' par élimination des coupures, alors Π et Π' ont même sémantique cohérente.*

DÉMONSTRATION: Il est clair qu'il suffit d'envisager le cas où Π' est obtenue par une étape élémentaire d'élimination des coupures.

On remarque alors que les expériences réussies de Π et Π' se correspondent bijectivement. $\diamond$

THÉORÈME 7.21. *Tout réseau ordonné admet une sémantique cohérente non triviale préservée par élimination des coupures.*

DÉMONSTRATION: Conséquence directe des lemmes précédents. $\diamond$

§D. Modularité de ce Calcul

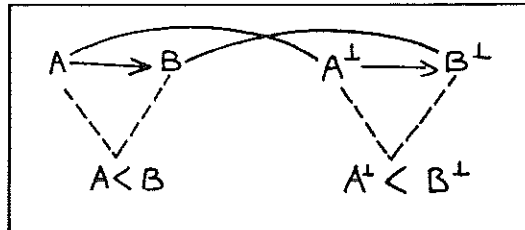
On exploite ici les résultats du paragraphe §G. du chapitre 1 pour définir des notions de module et de typage. Cela généralise fortement les notions de modules pour des réseaux commutatifs envisagées dans [DR90, Gir87b], en ce sens qu'on ne demande pas qu'un module soit un réseau avec hypothèses, et lorsqu'on branche deux modules on ne demande pas qu'ils soient disjoints: on prend simplement la réunion ensembliste de deux graphes. Un module est pour nous n'importe quelle partie d'un réseau, et il est possible qu'il ne contienne qu'une seule des deux arêtes d'un même lien. Par la même cela ressemble un peu à certaines constructions que [Mét92] a introduit pour les réseaux commutatifs.

DÉFINITION 7.22. *Un module M de bord Γ est un couple $\bar{\Pi}, \Gamma$ où $\bar{\Pi}$ est un sous-graphe d'un réseau ordonné, et Γ un sous ensemble de ses sommets. On définit les graphes de correction d'un module comme ceux d'un réseau: ce sont ses sous-graphes bigarrés maximaux; un parcourt du module associé à un graphe de correction G est alors la restriction à Γ du préordre u suivant: XuY si et seulement si G contient un chemin de X à Y . Soit M et M' deux modules de même bord Γ , et dont les sommets communs soient Γ et tel que si $M|_\Gamma$ contient l'arc $A \rightarrow B$ alors $M'|_\Gamma$ ne contient pas l'arc $B \rightarrow A$ et réciproquement; on définit alors leur branchement comme la réunion ensembliste de ces deux graphes si elle est un préréseau.*

Le théorème 1.33. du chapitre 1 montre que, le branchement de deux modules est un réseau correct si et seulement si toute restriction aux frontières des parcoures sont orthogonales. Il est alors naturel de définir le type d'un module de bord Γ comme un ensemble clos par bi-orthogonal de relations binaires sur Γ .

§E. η -expansion

L'axiome $A < B \text{---} A^\perp < B^\perp$ se décompose en



qu'on peut considérer comme un module de frontière $A < B, A^\perp < B^\perp$ dont le seul parcourt est $u = \{(A < B, A^\perp < B^\perp)\}$. Soit Π un réseau ordonné utilisant l'axiome $A < B \text{---} A^\perp < B^\perp$. Le module de frontière $A < B, A^\perp < B^\perp$ constitué par Π privé de cet axiome n'a que des parcours orthogonaux à ceux de l'axiome $A < B \text{---} A^\perp < B^\perp$ qui n'a qu'un seul parcourt, u . Le remplacement dans Π de $A < B \text{---} A^\perp < B^\perp$ par sa décomposition donne donc un réseau correct.

On démontrerait de même qu'un axiome $A \otimes B \text{---} A^\perp \otimes B^\perp$ peut être remplacé par sa décomposition.

Enfin une induction triviale montre que tout réseau ordonné peut être transformé en un réseau ordonné n'utilisant que des axiomes atomiques.

§F. Rapport avec le calcul de la deuxième partie.

PROPOSITION 7.23. *Un réseau ordonné sans coupure dont on oublie les arcs est un réseau habituel dans le quel les précédents sont devenus des par.*

DÉMONSTRATION: Il est clair qu'on obtient ainsi un préréseau. Il suffit alors de remarquer que tout chemin bigarré du réseau commutatif sous-jacent est un chemin bigarré du réseau ordonné. $\diamond$

THÉORÈME 7.24. *(plongement du calcul ordonné) Il est possible, dans un réseau ordonné, de transformer certains précédents en par et d'autres en tenseur et d'obtenir ainsi un réseau habituel.*

DÉMONSTRATION: On peut supposer sans perte de généralité que l'ordre sur les conclusions et coupures est l'ordre discret, et que le réseau est connexe.

On procède par induction sur le nombre de liens autres qu'axiomes du réseau Π .

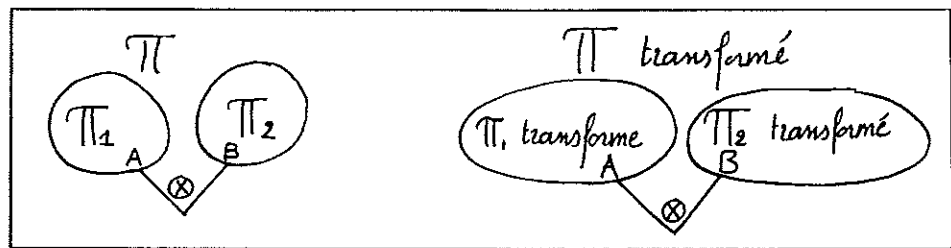
S'il n'y a que des axiomes, il n'y a rien à démontrer.

S'il y a un lien par ou précédent final, en le supprimant on le transforme en un réseau plus petit que l'on transforme; en remettant un lien par à la place du lien supprimé, on obtient une solution pour Π .

Sinon, le réseau commutatif sous-jacent contient

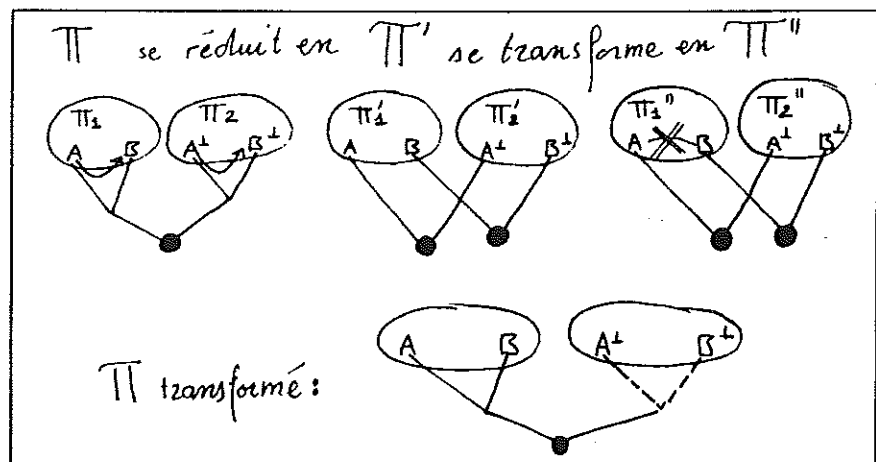
un tenseur ou une coupure héréditairement scindant d'après le lemme 4.30..

S'il s'agit d'un tenseur t , en le supprimant, on obtient deux réseaux plus petits, pour les quels on dispose d'une solution. En remettant ce tenseur supprimé on obtient une solution pour Π .

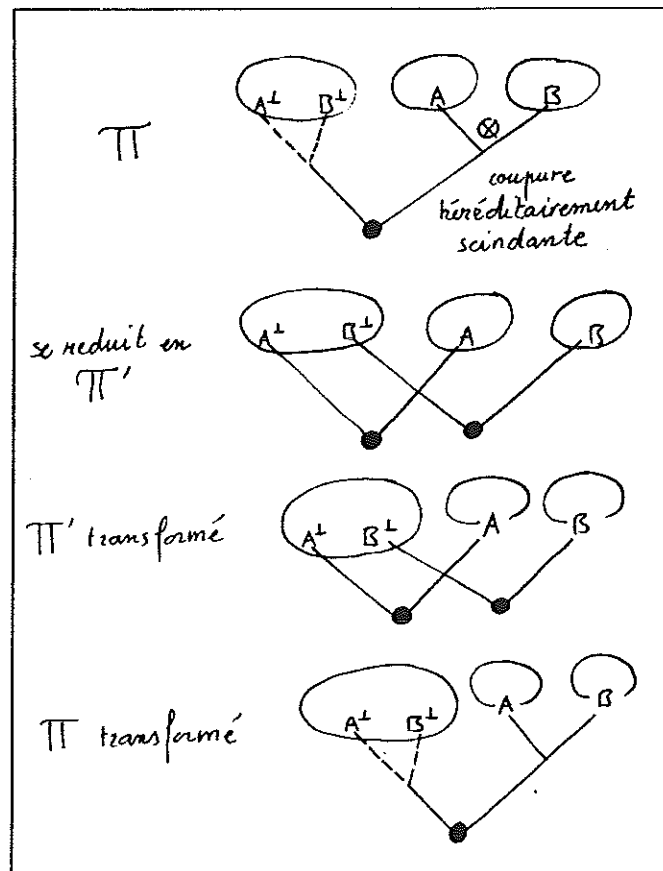


S'il s'agit d'une coupure, on envisage deux cas:

c'est une coupure entre deux précédés $A^\perp < B^\perp$ et $A < B$ On procède alors à une étape d'élimination des coupures sur cette coupure, et on supprime l'ordre apparu. Le réseau obtenu ayant un lien de moins, on dispose d'une solution Π' pour ce réseau. Comme la coupure réduite était scindante, les deux arêtes $A-A^\perp$ et $B-B^\perp$ constituent un ensemble séparateur de ce réseau Π' , appelons Π'_1 et Π'_2 ces deux morceaux — qui sont des réseaux commutatifs, dont les chemins sont à double sens. Si Π'_1 contenait un chemin bigarré reliant $A^\perp$ et $B^\perp$ et Π'_2 un chemin bigarré reliant A et B , alors Π' contiendrait un circuit bigarré empruntant les deux coupures et ces deux chemins, qui, étant dans des parties disjointes du réseau ne peuvent avoir de couleur commune. Le réseau obtenu en plaçant un lien **tenseur** sous les deux formules non-relies par un chemin bigarré, un lien **par** sous les deux autre, et une coupure entre ce **par** et ce **tenseur** est une solution pour Π .



c'est une coupure entre un par et un tenseur $A^\perp \otimes B^\perp$ et $A \otimes B$ On procède alors à une étape d'élimination des coupures sur cette coupure, et le réseau obtenu ayant un lien de moins, on dispose d'une solution Π' pour ce réseau. Comme la coupure réduite était scindante, les deux coupures $A-A^\perp$ et $B-B^\perp$ constituent un ensemble séparateur de ce réseau Π' , appelons Π'_1 et Π'_2 ces deux morceaux — qui sont des réseaux commutatifs, dont les chemins sont à double sens. Comme on a choisi une coupure *héréditairement* scindante le réseau Π'_2 ne peut contenir de chemin bigarré entre A et B ; la transformation ne crée pas de chemin (et donc a fortiori pas de chemin bigarré) entre des parties totalement déconnectées (ici la composante connexe de A et celle de B dans $\Pi'_2 - \{A \otimes B\}$). On obtient donc une solution pour Π en supprimant les deux coupures créées et en remettant les liens **tenseur**, **par** et coupure disparus.



◇

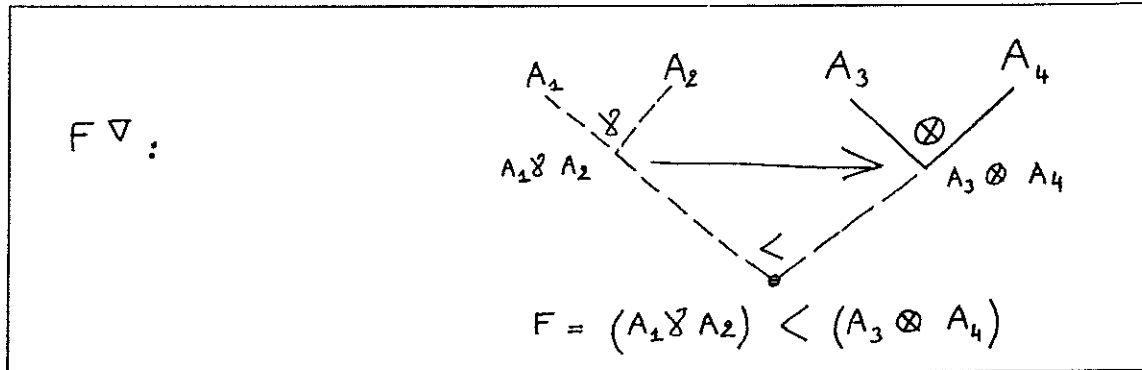
§G. Connecteurs n-aires définissables

Ce paragraphe fait grand usage du chapitre 1.

Dans le calcul commutatif tout préréseau Π de conclusions $A_1, \dots, A_n$ se transforme canoniquement en un préréseau Π' ayant une seule conclusion $F(A_1, \dots, A_n)$ de sorte que Π soit un réseau si et seulement si Π' en est un. On sait de plus que cette formule $F(A_1, \dots, A_n)$ est unique à l'associativité et la commutativité près du par et vaut $A_1 \wp \dots \wp A_n$.

Qu'en est-il pour ce calcul? On peut se demander ici quels sont les ordres i tels que tout préréseau de conclusion $A_1, \dots, A_n[i]$ se transforme de même en un préréseau Π' ayant une unique conclusion $F(A_1, \dots, A_n)$ où F est une formule écrite avec les A_i et nos connecteurs de sorte que Π soit un réseau si et seulement si Π' en est un.

Appelons F^∇ le module de bord $A_1, \dots, A_n$, constitué par l'arbre des sous-formules de F arrêté en $A_1, \dots, A_n$ (au sens de la définition 7.3., qui en fait un module), et notons $\mathcal{F}$ l'ensemble de ses parcours.



Notre question se ramène donc à la suivante:

existe-il $F(A_1, \dots, A_n)$ telle que, pour tout prémodule² M de bord $A_1, \dots, A_n$, M complété par i est un réseau, si et seulement si M complété par F^∇ est un réseau.

Nous pouvons maintenant formaliser cette question:

étant donné un ordre i sur $A_1, \dots, A_n$ existe-t-il F telle que u sur $A_1, \dots, A_n$ est orthogonale à tout parcourt de F^∇ si et seulement si u est orthogonale à i ?

Ce qui s'énonce plus concisément en utilisant l'orthogonalité sur les ensemble de relations:

existe-t-il F telle que $\mathcal{F}^\perp = \{i\}^\perp$?

PROPOSITION 7.25. *Si il existe $F(A_1, \dots, A_n)$ écrite avec les connecteurs $\wp$, $\otimes$ et $<$ telle que $\mathcal{F}^\perp = \{i\}^\perp$ alors F ne contient pas de connecteur $\otimes$.*

DÉMONSTRATION: En effet, si tel était le cas, soit A_i située de l'un des cotés du tenseur, et A_j de l'autre. Alors il existe un parcourt v de F^∇ contenant (A_i, A_j) et (A_j, A_i) . Si i ne contient pas (A_j, A_i) , alors $u = \{(A_i, A_j)\}$ est orthogonale à i sans être orthogonale à $v \in \mathcal{F}^\perp$. Même démonstration si i ne contient pas (A_i, A_j) . De plus, comme i est un ordre, i ne peut contenir (A_i, A_j) et (A_j, A_i) . $\diamond$

LEMME 7.26. *Soient $F_1, \dots, F_p$ des formules écrites avec les lettres $A_1, \dots, A_n$ de sorte que les lettres utilisées par $F_1, \dots, F_p$ constituent une partition de $A_1, \dots, A_p$. Soit j un ordre sur $F_1, \dots, F_p$ et j^l l'ordre obtenu par expansion des $F_1, \dots, F_p$ (on a vu au chapitre 1 paragraphe §B. comment s'obtient cet ordre). Soit M le module de bord $A_1, \dots, A_n$ défini par la réunion des $F_1^\nabla, \dots, F_p^\nabla$ et l'ordre j , et soit $\mathcal{M}$ l'ensemble de ses parcoures.*

Alors on a $\mathcal{M}^\perp = \{j^l\}^\perp$

DÉMONSTRATION: On procède par récurrence sur le nombre de liens de M . Si M ne contient pas de lien, alors $n = p$, $\forall i \leq n$ $A_i = F_i$ (à une permutation des indices près) et $j^l = j$ et le seul parcourt de M est $j = j^l$, et le résultat est évident.

Sinon, l'une des formules F_i est $G \wp H$ ou $G < H$. Remarquons que les lettres utilisées par les formules $G, H, F_2, \dots, F_n$ constituent aussi une partition des $A_1, \dots, A_n$. Munissons ces formules de l'ordre $\mathfrak{k} = j / (H \wp G) \rightsquigarrow H \sim G$ ou $\mathfrak{k} = j / (G < H) \rightsquigarrow G \lesssim H$ selon le cas. Appelons N le modules de bord $A_1, \dots, A_n$ constitué par la réunion des $G^\nabla, H^\nabla, F_2^\nabla, \dots, F_n^\nabla$ et l'ordre $\mathfrak{k}$, et

²un prémodule est module qui contient éventuellement des circuits bigarrés

appelons $\mathcal{N}$ l'ensemble de ses parcours. L'expansion de $\mathfrak{k}$ suivant les formules $G, H, F_2, \dots, F_n$ est précisément j^1 , et en vertu de l'hypothèse de récurrence on a donc $\mathcal{N}^\perp = \{j^1\}^\perp$. Il suffit donc de montrer que $\mathcal{M}^\perp = \mathcal{N}^\perp$ pour pouvoir conclure. On va pour cela montrer que, pour toute relation u sur $A_1, \dots, A_n$ on a

$$(\exists v \in \mathcal{M} \ \neg(u \perp v)) \iff (\exists v' \in \mathcal{N} \ \neg(u \perp v'))$$

(i) Supposons que $F_1 = G \mathfrak{H} H$.

(i.a) (Implication directe) Supposons qu'il existe un circuit bigarré dans $M \cup u$, et montrons qu'il en existe un dans $N \cup u$. Pour cela, prenons un circuit bigarré c de $M \cup u$ de longueur minimale. Remarquons que c , qui est bigarré, n'utilise pas les deux arêtes $H \text{---} G \mathfrak{H} H$ et $G \text{---} G \mathfrak{H} H$. Si c n'utilise aucune des arêtes $H \text{---} G \mathfrak{H} H$ et $G \text{---} G \mathfrak{H} H$ alors ce circuit bigarré est un circuit bigarré de $N \cup u$. S'il en utilise une, disons $H \text{---} G \mathfrak{H} H$, l'arc suivant de c est $H \mathfrak{H} G \rightarrow X$ (respectivement $H \mathfrak{H} G \leftarrow X$). Par définition de l'ordre $\mathfrak{k} = j / (G \mathfrak{H} H) \sim G \sim H$, $N \cup u$ contient le circuit bigarré c' obtenu par remplacement de $H \text{---} G \mathfrak{H} H \rightarrow X$ (respectivement $H \text{---} G \mathfrak{H} H \leftarrow X$) par $H \rightarrow X$ (respectivement $H \leftarrow X$).

(i.b) (Implication réciproque) Supposons qu'il existe un circuit bigarré de $N \cup u$, et montrons qu'il en existe un dans $M \cup u$. Pour cela, prenons un circuit bigarré c de $M \cup u$ de longueur minimale. Si c n'emprunte pas d'arc incident à H ou G , c'est un circuit bigarré de $M \cup u$. Remarquons qu'un circuit passant par H (ou par G) emprunte un arc de $\mathfrak{k}$ incident à H (ou G), puisque les seules arêtes incidentes à H sont de même couleur.

Remarquons ensuite que ce circuit bigarré étant de longueur minimale, s'il emprunte un arc incident à G (resp H), alors il ne passe pas par H (resp G): en effet si tel était le cas, comme H et G sont équivalents dans $\mathfrak{k}$, $G \dots s \dots H \rightarrow X \dots t \dots G$ pourrait être remplacée par $G \rightarrow X \dots t \dots G$, et $G \dots s \dots H \leftarrow X \dots t \dots G$ pourrait être remplacée par $G \leftarrow X \dots t \dots G$.

On est donc ramené, à la symétrie des rôles de H et G près au cas où ce circuit passe par H et pas par G et emprunte un arc incident à H . Il suffit alors de remplacer l'arc $H \rightarrow X$ (ou $H \leftarrow X$) par $H \text{---} H \mathfrak{H} G \rightarrow X$ (ou par $H \text{---} H \mathfrak{H} G \leftarrow X$) pour obtenir un circuit bigarré de $M \cup u$. Le fait qu'on n'ait pas à remplacer simultanément un arc incident à G par une telle séquence nous assure que le circuit ainsi constitué ne contient pas les deux arêtes $H \text{---} G \mathfrak{H} H$ et $G \text{---} G \mathfrak{H} H$, et est donc bigarré).

(ii) Supposons que $F_1 = G < H$.

(ii.a) (Implication directe) On procède comme en (i.a) si ce n'est que le circuit c peut emprunter l'arc du lien $G < H$. Mais on remarque que c'est aussi un arc de $\mathfrak{k}$, et c est un circuit de $N \cup u$.

(ii.b) (Implication réciproque) On procède comme en (i.b) si ce n'est que le circuit peut passer et par H et par G et être de longueur minimale. Dans ce cas il emprunte l'arc $G \rightarrow H$ de l'ordre $\mathfrak{k}$. C'est alors un circuit bigarré de $M \cup u$ empruntant l'arc $G \rightarrow H$ du lien $G < H$.

Appliquons maintenant ce lemme dans le cas où $p = 1$:

PROPOSITION 7.27. *L'ensemble des relations orthogonales aux parcours $\mathcal{F}$ de F^∇ est l'ensemble des relations orthogonales à j^1 , l'ordre obtenu par expansion de F .*

Notre question se ramène donc à:

quand existe-t-il une formule F , qu'on peut supposée écrite avec les connecteurs $\wp$ et $<$ tel que $i = j^1$, où j^1 est l'ordre expansé correspondant à F ?

Nous y avons déjà répondu au chapitre 1 par le théorème 1.18.:

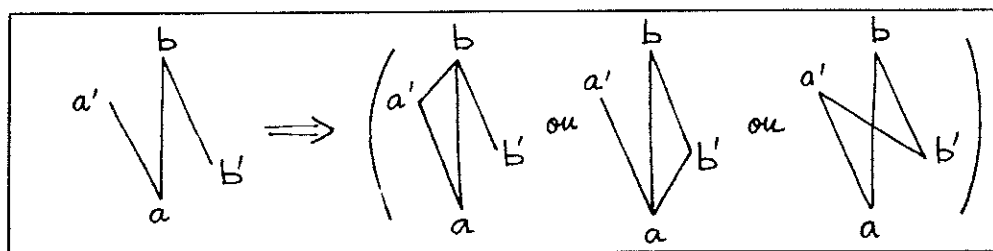
THÉORÈME 7.28. *Il existe une formule $F(A_1, \dots, A_n)$ équivalente à $A_1, \dots, A_n[i]$ i.e. il existe F telle que les parcours $\mathcal{F}$ de F^∇ satisfasse:*

$$\forall u (u \perp \mathcal{F} \iff u \perp i)$$

si et seulement si i est contractile, et dans ce cas F est unique à l'associativité de $\wp$ et $<$ près, et à la commutativité de $\wp$ près.

Rappelons que les ordres contractiles, qui sont les expansés de l'ordre trivial, ont été caractérisés dans le théorème 1.13. par la condition globale $\mathbb{P}$:

$$\mathbb{P} : \quad \forall a, b, a', b' \begin{pmatrix} a < a' \\ \text{et} \\ a < b \\ \text{et} \\ b' < b \end{pmatrix} \Rightarrow \begin{pmatrix} a' \leq b \\ \text{ou} \\ a \leq b' \\ \text{ou} \\ b' \leq a' \end{pmatrix}$$



Chapitre 8

Le Calcul des Séquents Ordonnés

§A. Présentation

Les règles du calcul des séquents correspondant doivent opérer et sur un multi-ensemble ordonné de formules: les connecteurs **par** et **tenseur** ayant la même sémantique et les mêmes liens que dans le calcul habituel, nous demanderons que les règles **TENSEUR** et **PAR** opèrent sur les formules comme dans le calcul usuel. Quant au connecteur **précède**, vu qu'il est situé entre le **par** et le **tenseur**, on peut à juste titre se demander s'il lui faut une règle à deux prémisses ou/et une règle à une prémisse. En fait nous allons lui trouver une règle à une prémisse, puis démontrer que toute règle correcte (dans les réseaux) à deux prémisses est simulable dans notre calcul. Par simulable nous voulons dire que:

tout résultat de l'application d'une éventuelle règle **précède** (correcte) à deux séquents-prémisses, peut s'obtenir par application de deux de nos règles à ces mêmes séquents prémisses

Cette notion est plus faible que dérivable, car on n'a pas démontré que le remplacement d'une application de cette éventuelle règle à deux prémisses par une application de deux de nos règles était uniforme sur les ordres (sur les formules, il l'est évidemment), bien que cela soit vraisemblable.

Cela montre qu'il n'est pas nécessaire de lui trouver une règle à deux prémisses. Notons d'ailleurs que dans les réseaux le comportement du **précède** est bien plus proche de celui du **par** que de celui du **tenseur**.

Il ne reste donc plus qu'à préciser la manière dont ces règles agissent sur l'ordre.

Cette action se devra de respecter les principes suivants:

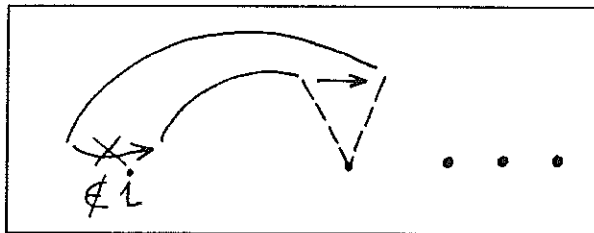
- (i) Les règles sont interprétées par la sémantique cohérente.
- (ii) Les règles ne donnent naissance qu'à des réseaux ordonnés corrects.
- (iii) Les règles permettent de séquentialiser le plus de réseaux corrects possible.

On remarquera que, pour que les deux premiers principes soient satisfaits, il faut que l'ordre sur le séquent conclusion ne soit pas trop grand (pour l'inclusion); par contre, le troisième principe nous incite à le prendre le plus grand possible. On dira qu'une règle est plus libérale qu'une autre si l'ordre sur le séquent conclusion est plus grand, i.e. pour tout application de la règle la moins libérale il est une application de la règle plus libérale qui fournisse un ordre plus grand.

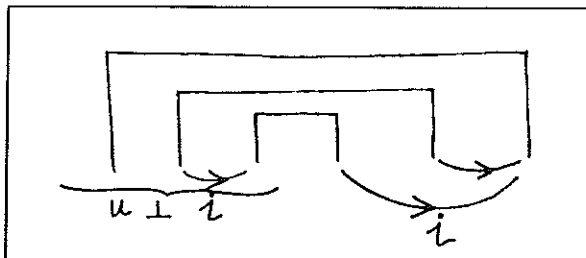
La mise au point de ce calcul utilise de manière essentielle les résultats du chapitre 1.

Nous procéderons pour chaque règle ainsi: on détermine la règle la plus libérale qui ne crée que des réseaux corrects, et on montre qu'elle est interprétée par la sémantique cohérente, ce qui suffit. On pourrait même montrer que la règle la plus libérale qui soit interprétée par la sémantique cohérente est la même que la règle la plus libérale qui ne crée que des réseaux corrects. En effet le détail des preuves de validité pour la sémantique cohérente utilise crucialement toutes les hypothèses signifiant que le réseau est correct. Cela n'est pas trop étonnant, compte tenu du théorème 7.21. qui montre que tout réseau correct admet un sémantique cohérente non-triviale.

Pour déterminer la règle la plus libérale à ne créer que des réseaux corrects, le lecteur aura peut-être l'impression qu'on utilise l'argument suivant: l'ensemble des parcours des réseaux de conclusions $A_1, \dots, A_n[i]$ construit par des preuves séquentielles de $\vdash A_1, \dots, A_n[i]$ contient toutes les relations orthogonales à i . En fait cet argument n'est qu'un abus de langage, car un examen méticuleux des démonstrations où semble intervenir cet argument montre qu'on n'utilise que l'assertion bien plus faible suivante (utilisée dans le lemme 7.26.: si i ne contient pas $A_i < A_j[i]$ alors il existe une preuve séquentielle dont le réseau associé contient un parcours contenant (A_j, A_i)). Justifions la rapidement. Le réseau suivant, qui est séquentialisable avec toutes les règles imaginables, montre que dès qu'il y a plus de deux formules c'est le cas.



Ensuite on remarquera que des règles raisonnables ne devrait pas faire intervenir le nombre de formules dans un séquent prémisses. Finalement — et c'est là l'argument qui motive notre choix — on doit pouvoir considérer des preuves avec des séquents-hypothèses, qui dans les réseaux peuvent se traduire par et l'argument le plus fort est alors pleinement justifié.



§B. Règles à une prémisses

¶B.1. L'axiome

$$\overline{\vdash A, A^\perp [\emptyset]} \text{ AXIOME}$$

L'ordre vide est assurément le plus grand qu'on puisse prendre pour obtenir un graphe simple... Le réseau ainsi constitué est correct, et la règle est trivialement interprétée par la sémantique cohérente.

DÉMONSTRATION: Soient deux n-uplets cohérents

$$(a_\bullet, a_\circ, a_1, \dots, a_n) \odot (a'_\bullet, a'_\circ, a'_1, \dots, a'_n) \left[\prod_i A_\bullet, A_\circ, A_1, \dots, A_n \right]$$

montrons qu'on a

$$((a_\bullet, a_\circ), a_1, \dots, a_n) \odot ((a'_\bullet, a'_\circ), a'_1, \dots, a'_n) \left[\prod_j (A_\bullet \wp A_\circ), A_1, \dots, A_n \right] \quad \text{où } j = i/\sim_\bullet \sim_\circ \wp_\bullet$$

Si i est un indice (éventuellement $\bullet$ ou $\circ$), l'expression $i : \wedge$ signifiera $a_i \wedge a'_i$

Si les deux premiers n-uplets sont égaux, les deux seconds aussi.

Si $(a_\bullet, a_\circ, a_1, \dots, a_n) \wedge (a'_\bullet, a'_\circ, a'_1, \dots, a'_n)$, on va envisager trois cas, suivant l'indice $\bullet$, $\circ$ ou i qui les rend cohérents.

Si $\bullet : \wedge$ et $\forall i > \bullet[i] \ i :=$, on a $(\bullet, \circ) : \wedge$; de plus, par construction de j on a $i > (\bullet, \circ)[j]$ si et seulement si $i > \bullet[i]$ et donc $i :=$.

Le cas $\circ : \wedge$ et $\forall i > \circ \ i :=$, est symétrique.

Si $i : \wedge$ et $\forall i' > i[i] \ i' :=$ alors i convient: si $(\bullet, \circ) > i[i]$ alors $\bullet > i[i]$ et $\circ > i[i]$ et donc $(\bullet, \circ) :=$. Si $i' > i[i]$ alors $i' > i[i]$ et donc $i' :=$. $\diamond$

¶B.4. Le connecteur précède

Nous déterminons ici la règle du prémisses à une prémisses et nous montrerons à la fin du chapitre qu'elle permet, avec la règle de MÉLANGE donnée plus tard de simuler toute règle correcte à deux prémisses pour ce connecteur.

Le connecteur **précède** correspond à la contraction de deux formules équivalentes-inférieures dans l'ordre

$$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \\ \vdots \end{array} \quad \vdash A_\bullet, A_\circ, A_1, A_2, \dots, A_n \ [i]}{\vdash (A_\bullet < A_\circ), A_1, A_2, \dots, A_n \left[i /_{A_\bullet \lesssim A_\circ \rightsquigarrow (A_\bullet < A_\circ)} \right]} \text{PRÉCÈDE}$$

PROPOSITION 8.3. Cette règle est la plus libérale à ne créer que des réseaux corrects.

DÉMONSTRATION: Même démonstration que pour la règle du par . $\diamond$

4.a. Sémantique cohérente de la règle du précède

PROPOSITION 8.4. Cette règle est interprétée par la sémantique cohérente.

DÉMONSTRATION: Soient deux n-uplets cohérents

$$(a_\bullet, a_\circ, a_1, \dots, a_n) \odot (a'_\bullet, a'_\circ, a'_1, \dots, a'_n) \left[\prod_j A_\bullet, A_\circ, A_1, \dots, A_n \right]$$

montrons qu'on a

$$((a_\bullet, a_\circ), a_1, \dots, a_n) \odot ((a'_\bullet, a'_\circ), a'_1, \dots, a'_n) \left[\prod_j (A_\bullet < A_\circ), A_1, \dots, A_n \right] \quad \text{où } j = i /_{\bullet, \circ, \dots, \mathfrak{F}_\circ}$$

Si les deux premiers n-uplets sont égaux, les deux seconds aussi.

Si $(a_\bullet, a_\circ, a_1, \dots, a_n) \wedge (a'_\bullet, a'_\circ, a'_1, \dots, a'_n)$, on va envisager trois cas, suivant l'indice $\bullet$, $\circ$ ou i qui les rend cohérents.

Si $\bullet : \wedge$ et $\forall i > \bullet[i] \ i :=$, comme $\circ :=$ on a $(\bullet, \circ) : \wedge$; de plus, par construction de j on a $i > (\bullet, \circ)[j]$ si et seulement si $i > \bullet[i]$ et donc $i :=$.

Si $\circ : \wedge$ et $\forall i > \circ \ i :=$, alors $(\bullet, \circ) : \wedge$, et comme $i > (\bullet, \circ)[j]$ si et seulement si $i > \bullet[i]$ on a $i :=$.

Si $i : \wedge$ et $\forall i' > i[i] \ i' :=$ alors i convient: si $(\bullet, \circ) > i[i]$ alors $\bullet > i[i]$ et $\circ > i[i]$ et donc $(\bullet, \circ) :=$. Si $i' > i[i]$ alors $i' > i[i]$ et donc $i' :=$. $\diamond$

§C. Règles à deux prémisses

L'ordre précédemment obtenu pour les règles à une prémisses était spécialement simple à décrire, tandis que pour les règles à deux prémisses, il faudra faire preuve de plus de circonspection, c'est à dire utiliser les résultat du chapitre 1 paragraphe §F..

¶C.1. La règle de MÉLANGE

L'application de la règle de MÉLANGE à deux preuves séquentielles π_1 et π_2 s'écrit:

$$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdash A_1, A_2, \dots, A_n [i] \end{array} \quad \begin{array}{c} \vdots \pi_2 \\ \vdash B_1, B_2, \dots, B_p [j] \end{array}}{\vdash A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p [\mathfrak{E}]} \text{ MÉLANGE}$$

Appelons π la preuve séquentielle ainsi obtenue. Exprimons maintenant le fait que cette règle ne crée que des réseaux corrects à l'aide de la proposition 7.11.:

- * Le réseau Π^- obtenu en omettant l'ordre sur les conclusion est correct (sans circuit bigarré).
- * L'ordre sur les conclusions est orthogonal à tout parcourt du réseau.

Appelons Π_1 le réseau correspondant à π_1 et Π_2 le réseau correspondant à π_2 . Les parcoures de Π^- sont les réunions des parcoures de Π_1^- et Π_2^- , de domaines disjoints. Nous avons donc déjà étudié au chapitre 1 paragraphe §F., à quelle condition $\mathfrak{E}$ est orthogonal à toute telle réunion de deux relations dont les domaines sont disjoints et nous avons appelés ces ordres des extensions orthogonales de i, j . La règle la plus libérale qu'on puisse prendre est donc assurément la suivante:

$$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \end{array} \quad \frac{\vdash A_1, A_2, \dots, A_n [i] \quad \vdash B_1, B_2, \dots, B_p [j]}{\vdash A_1, A_2, \dots, A_n, B_1, B_2, \dots, B_p [\ell]} \text{MÉLANGE}}$$

où ℓ est une extension orthogonale de i, j .

Notons qu'il existe toujours de telles extensions orthogonales, telles $\ell = i \cup j$, $\ell = i < j$, $\ell = j < i$...

Rappelons qu'on a caractérisé (théorème 1.27.) ces ordres par :

ℓ est une extension orthogonale de deux ordres i et j définis sur deux ensembles disjoints E et F si et seulement si:

$$(i) \quad \ell|_E = i$$

$$(ii) \quad \ell|_F = j$$

$$(iii) \quad T(\ell, E, F) \text{ avec } T(\ell, E, F): \forall x, y \in E \quad \forall x', y' \in F \quad \left| \begin{array}{c} x \ell x' \\ \text{et} \\ y' \ell y \end{array} \right| \Rightarrow \left| \begin{array}{c} x \ell y \\ \text{ou} \\ y' \ell x' \end{array} \right|$$

Cela nous sera très utile dans le paragraphe suivant:

1.a. Sémantique cohérente de la règle de MÉLANGE

PROPOSITION 8.5. Cette règle est interprétée par la sémantique cohérente.

DÉMONSTRATION: On suppose que

$$(a_1, a_2, \dots, a_n) \odot (a'_1, a'_2, \dots, a'_n) \left[\prod_i A_i \right]$$

et que

$$(a_{n+1}, \dots, a_m) \odot (a'_{n+1}, \dots, a'_m) \left[\prod_j A_j \right]$$

on va montrer que

$$(a_1, a_2, \dots, a_n, a_{n+1}, \dots, a_m) \odot (a'_1, a'_2, \dots, a'_n, a'_{n+1}, \dots, a'_m) \left[\prod_{\ell} A_k \right]$$

On écrira $l : \odot$ pour $a_l \odot a'_l$.

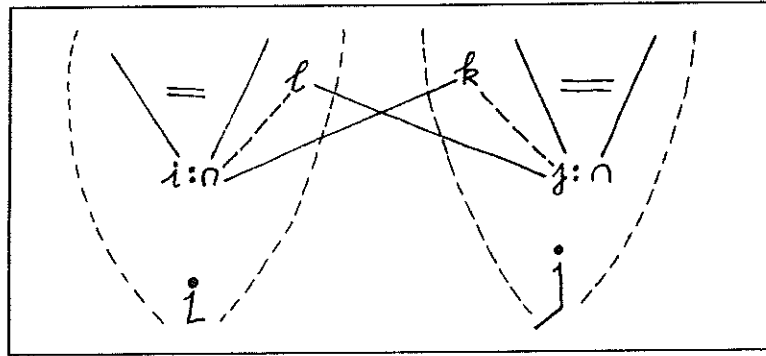
Les lettres i, i' désigneront des indices $\leq n$ (correspondant à des éléments de i), et les lettres j, j' des indices $> n$ (correspondant à des éléments de j). On est nécessairement dans l'un des quatre cas suivants:

Si $(a_1, a_2, \dots, a_n) = (a'_1, a'_2, \dots, a'_n)$ et $(a_{n+1}, \dots, a_m) = (a'_{n+1}, \dots, a'_m)$ alors le résultat est évident.

Si $(a_1, a_2, \dots, a_n) = (a'_1, a'_2, \dots, a'_n)$ et $(a_{n+1}, \dots, a_m) \odot (a'_{n+1}, \dots, a'_m)$, alors il existe un indice $j : \sim$ tel que tous les $j' > j[\mathfrak{E}]$ soient $j :=$. On va montrer que cet indice a la même propriété dans $\mathfrak{E}$. Prenons un indice qui lui soit supérieur dans $\mathfrak{E}$. Si c'est un indice de j , disons j' , comme $\mathfrak{E}$ étend conservativement j (ii), cet indice lui était supérieur dans j , et donc $j' :=$. Si c'est un indice de i , on a de toute façon $i :=$.

Le cas $(a_1, a_2, \dots, a_n) \odot (a'_1, a'_2, \dots, a'_n)$ et $(a_{n+1}, \dots, a_m) = (a'_{n+1}, \dots, a'_m)$ est symétrique.

Si $(a_1, a_2, \dots, a_n) \wedge (a'_1, a'_2, \dots, a'_n) [\prod_i A_i]$ et $(a_{n+1}, \dots, a_m) \wedge (a'_{n+1}, \dots, a'_m) [\prod_j A_j]$, alors on a un indice $i : \sim$ de i tel que tout indice i' de i situé au dessus de lui dans i soit $i' :=$, et aussi un indice $j : \sim$ de j tel que tout indice j' de j situé au dessus de lui dans j soit $j' :=$. Si ni l'un ni l'autre ne satisfaisait cette même propriété dans $\mathfrak{E}$ tout entier, alors il existerait un indice k au dessus de i dans $\mathfrak{E}$ tel que $k \neq$ et un indice l au dessus de j' dans $\mathfrak{E}$ tel que $l \neq$. Comme $\mathfrak{E}$ étend conservativement i (i) et j (ii), k est un point de j et l un point de i ; comme $\neg(k > j[\mathfrak{E}])$ et $\neg(l > i[\mathfrak{E}])$, d'où $\neg(k > j[j])$ et $\neg(l > i[i])$, ceci contredirait (iii). L'un des deux indices $i : \sim$ et $j : \sim$ n'a donc que des majorants $h :=$ dans $\mathfrak{E}$.



◇

¶C.2. Le TENSEUR ordonné

L'application de la règle de TENSEUR à deux preuves séquentielles π_1 et π_2 s'écrit:

$$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdash A_*, A_1, A_2, \dots, A_n [i] \end{array} \quad \begin{array}{c} \vdots \pi_2 \\ \vdash A_*, A_{n+1}, \dots, A_m [j] \end{array}}{\vdash (A_* \otimes A_*) A_1, A_2, \dots, A_n, A_{n+1}, \dots, A_m [\mathfrak{E}]} \text{TENSEUR}$$

Appelons π la preuve séquentielle ainsi obtenue. Exprimons maintenant le fait que cette règle ne crée que des réseaux corrects à l'aide de la proposition 7.11.:

- * Le réseau Π^- obtenu en omettant l'ordre sur les conclusion est correct (sans circuit bigarré).
- * L'ordre sur les conclusions est orthogonal à tout parcours du réseau.

Appelons Π_1 le réseau correspondant à π_1 et Π_2 le réseau correspondant à π_2 . Comme ce tenseur est final, chacune de ses deux arêtes, qui sont noires, est un isthme de Π^- . Les parcours

de Π^- sont donc les clôtures transitives des réunions des parcours de Π_1^- et Π_2^- , dans les quels on renomme $A_\bullet \otimes A_\circ$ les deux points $A_\bullet$ et $A_\circ$. Nous avons donc déjà étudié au chapitre 1 paragraphe §F., à quelle condition $\mathfrak{E}$ est orthogonal à toute telle réunion (l'orthogonalité est stable par clôture transitive) de deux relations dont la seule chose que l'on sache est qu'elles sont respectivement orthogonales à i et j ; et nous avons appelés ces ordres des extensions orthogonales de i, j au dessus du point commun, ici $A_\bullet \otimes A_\circ$. La règle la plus libérale qu'on puisse prendre est donc la suivante:

$$\frac{\vdots \pi_1 \quad \vdots \pi_2}{\vdash A_o, A_1, A_2, \dots, A_n [i] \quad \vdash A_o, A_{n+1}, \dots, A_m [j]} \text{ TENSEUR}$$

où ℓ est n'importe quelle extension orthogonale de i, j dans lesquels on renomme $A_\bullet \otimes A_0$ les points $A_\bullet$ et A_0^1

Notons qu'il existe toujours de telles extensions orthogonales, par exemple la relation obtenue en considérant i comme une relation sur $A_0 \otimes A_0, A_1, \dots, A_n$ et j comme une relation sur $A_0 \otimes A_0, A_{n+1}, \dots, A_m$, et en prenant la clôture transitive de la réunion de ces deux relations.

Rappelons qu'on a caractérisé (théorème 1.29. ces extensions orthogonales, et que compte tenu de l'identification à la quelle on procède, la caractérisation est la suivante, où $E = \{A_*, A_1, \dots, A_n\}$ et $F = \{A_0, A_{n+1}, \dots, A_m\}$:

$$\begin{aligned} (i): \quad \mathfrak{E}_{A_0 \otimes A_0 \rightsquigarrow A}|_E = \mathbf{i} & \qquad (i'): \quad \mathfrak{E}_{A_0 \otimes A_0 \rightsquigarrow A_0}|_F = \mathbf{i}' \\ (ii): \quad \mathbb{T}(\mathfrak{E}, E, F - \{A_0\}) & \qquad (ii'): \quad \mathbb{T}(\mathfrak{E}, E - \{A_0\}, F) \\ (iii): \quad \mathbb{T}_{A_0 \otimes A_0}(\mathfrak{E}, E - \{A_0\}, F - \{A_0\}) \end{aligned}$$

avec $\mathbb{T}_{A_\bullet \otimes A_\circ}(\mathfrak{k}, E - \{A_\bullet\}, F - \{A_\circ\})$:

$$\forall y \in E - \{A_\bullet\} \quad \forall y' \in F - \{A_\circ\} \quad \left| \begin{array}{l} y \mathfrak{E} y' \Rightarrow (y \mathfrak{E} A_\bullet \otimes A_\circ \text{ ou } A_\bullet \otimes A_\circ \mathfrak{E} y') \\ \text{et} \\ y' \mathfrak{E} y \Rightarrow (y' \mathfrak{E} A_\bullet \otimes A_\circ \text{ ou } A_\bullet \otimes A_\circ \mathfrak{E} y) \end{array} \right.$$

cette caractérisation est nécessaire au paragraphe suivant:

2.a. Sémantique cohérente de la règle du TENSEUR

PROPOSITION 8.6. *Cette règle est interprétée par la sémantique cohérente.*

DÉMONSTRATION: On suppose que

$$(a_{\bullet}, a_1, a_2, \dots, a_n) \odot (a'_{\bullet}, a'_1, a'_2, \dots, a'_n) \left[\prod_i A_i \right]$$

¹cette identification de deux points n'est pas très élégante, mais elle se comprend bien; une définition plus formelle eut exigé de faire des sommes amalgamées de relations, ce qui nous a semblé pire...

et que

$$(a_o, a_{n+1}, \dots, a_m) \odot (a'_o, a'_{n+1}, \dots, a'_m) \left[\prod_j A_j \right]$$

on va montrer que

$$((a_\bullet, a_o), a_1, a_2, \dots, a_n, a_{n+1}, \dots, a_m) \odot ((a'_\bullet, a'_o), a'_1, a'_2, \dots, a'_n, a'_{n+1}, \dots, a'_m) \left[\prod_{\mathfrak{k}} A_k \right]$$

On écrira $l : \odot$ pour $a_l \odot a'_l$.

Les lettres i, i' désigneront des indices $\leq n$ ou l'indice $\bullet$ (correspondant à des éléments de $|i|$), et les lettres j, j' des indices $> n$ ou l'indice $\circ$ (correspondant à des éléments de $|j|$).

Si $(a_\bullet, a_1, a_2, \dots, a_n) = (a'_\bullet, a'_1, a'_2, \dots, a'_n)$ et $(a_o, a_{n+1}, \dots, a_m) = (a'_o, a'_{n+1}, \dots, a'_m)$ alors le résultat est évident.

Si $(a_\bullet, a_1, a_2, \dots, a_n) = (a'_\bullet, a'_1, a'_2, \dots, a'_n)$ et $(a_o, a_{n+1}, \dots, a_m) \wedge (a'_o, a'_{n+1}, \dots, a'_m)$, appelons j l'indice de $\mathfrak{j}$ qui rendent les deux n-uplets strictement cohérents: $j : \wedge$ et pour tout j' de $\mathfrak{j}$ au dessus de j dans $\mathfrak{j}$ on a $j' := (j \text{ et } j' \text{ peuvent éventuellement valoir } \circ)$.

Si cet indice est $j = \circ$, montrons que $(\bullet, \circ)$ rend les deux n-uplets de $\mathfrak{k}$ cohérents. On a déjà $(\bullet, \circ) : \wedge$ (par la définition de la cohérence suivant le tenseur). Soit i un point de $\mathfrak{i}$ autre que $\bullet$ situé au dessus de $(\bullet, \circ)$ dans $\mathfrak{k}$; on a $i :=$. Soit maintenant j un point de $\mathfrak{j}$ autre que $\circ$ situé au dessus de $(\bullet, \circ)$ dans $\mathfrak{k}$; d'après la propriété (i') de $\mathfrak{k}$ ce point j est au dessus de $\circ$ dans $\mathfrak{j}$ et donc $j :=$.

Si cet indice est $j \neq \circ$, montrons que ce même indice convient dans $\mathfrak{k}$. On a déjà $j : \wedge$. Supposons que $(\bullet, \circ)$ soit au dessus de j dans $\mathfrak{k}$; d'après la propriété (i') de $\mathfrak{k}$, on a $\circ > j[j]$ et donc $\circ :=$ comme $\bullet :=$ on a $(\bullet, \circ) :=$. Soit maintenant $j' \neq \circ$ de $\mathfrak{j}$ au dessus de j dans $\mathfrak{k}$; d'après la propriété (i') de $\mathfrak{k}$, cet indice j' est au dessus de j dans $\mathfrak{j}$, et donc $j :=$. Soit maintenant i de $\mathfrak{i} - \bullet$ au dessus de j' dans $\mathfrak{k}$, on a $i' :=$.

Le cas $(a_\bullet, a_1, a_2, \dots, a_n) \wedge (a'_\bullet, a'_1, a'_2, \dots, a'_n)$ et $(a_o, a_{n+1}, \dots, a_m) = (a'_o, a'_{n+1}, \dots, a'_m)$ est symétrique.

Si $(a_\bullet, a_1, a_2, \dots, a_n) \wedge (a'_\bullet, a'_1, a'_2, \dots, a'_n) \left[\prod_i A_i \right]$ et $(a_o, a_{n+1}, \dots, a_m) \wedge (a'_o, a'_{n+1}, \dots, a'_m) \left[\prod_j A_j \right]$

on a un indice i de $\mathfrak{i}$ qui rend les deux premiers n-uplets strictement cohérents et un indice j de $\mathfrak{j}$ qui rend les deux seconds n-uplets cohérents. On va envisager quatre cas, suivant que ces indices font ou non partie de $\bullet, \circ$.

Si $i = \bullet$ et $j = \circ$, le point $(\bullet, \circ)$ de $\mathfrak{k}$ convient. On a déjà $(\bullet, \circ) : \wedge$. Soit maintenant un point $i \neq \bullet$ de $\mathfrak{i}$ situé dans $\mathfrak{k}$ au dessus de $(\bullet, \circ)$. La propriété (i) de $\mathfrak{k}$ montre que i est au dessus de $\bullet$ dans $\mathfrak{i}$ et donc $i :=$. Soit maintenant un point $j \neq \circ$ de $\mathfrak{j}$ situé dans $\mathfrak{k}$ au dessus de $(\bullet, \circ)$. La propriété (i') de $\mathfrak{k}$ montre que j est au dessus de $\circ$ dans $\mathfrak{j}$ et donc $j :=$.

Si $i = \bullet$ et $j \neq \circ$.

Si $\circ$ est au dessus de j dans $\mathfrak{j}$, (et donc $\circ :=$), alors $(\bullet, \circ)$ convient. On a déjà $(\bullet, \circ) : \wedge$ puisque $\bullet : \wedge$ et $\circ :=$. Soit maintenant $i \neq \bullet$ un point de $\mathfrak{i}$ situé au dessus de $(\bullet, \circ)$

dans $\mathfrak{k}$. D'après la propriété (i) de $\mathfrak{k}$, ce point i est au dessus de $\bullet$ dans $\mathfrak{i}$ et donc $i :=$. Soit maintenant un point j' au dessus de $(\bullet, o)$ dans $\mathfrak{k}$; d'après la propriété (i') de $\mathfrak{k}$ ce point j' est au dessus de o dans $\mathfrak{j}$; comme o est au dessus de j , ce point j' est au dessus de j dans $\mathfrak{j}$, et donc $j' :=$.

Si o n'est pas au dessus de j dans $\mathfrak{j}$, montrons que j convient. On a déjà $j : \wedge$.

Comme o n'est pas au dessus de j dans $\mathfrak{j}$ le point $(\bullet, o)$ n'est pas au dessus de j dans $\mathfrak{k}$ en vertu de la propriété (i') de $\mathfrak{k}$. Soit i un point de $\mathfrak{i} - \bullet$ situé au dessus de j dans $\mathfrak{k}$. En raison de la propriété (iii) de $\mathfrak{k}$, comme o n'est pas au dessus de j dans $\mathfrak{j}$ ce point est au dessus de $\bullet$ dans $\mathfrak{i}$ et donc $i :=$.

Si $i \neq \bullet$ et $j = o$. Sous-cas symétrique au précédent.

Si $i \neq \bullet$ et $j \neq o$, montrons que i ou j convient dans $\mathfrak{k}$. Si ce n'était pas le cas, il existerait j' de $\mathfrak{j}$ au dessus de i dans $\mathfrak{k}$ tel que $j' \neq$, et un i' de $\mathfrak{i}$ au dessus de j dans $\mathfrak{k}$ tel que $i' \neq$. Les propriétés (i) et (ii) de $\mathfrak{k}$ montre que $i' \neq \bullet$ et $j' \neq o$.

La propriété (ii) (ou (ii')) de $\mathfrak{k}$ montre que $i' > i[\mathfrak{k}]$ ou $j' > j[\mathfrak{k}]$ ce qui compte tenu de (i) (ii) et de $j' \neq$ et $i' \neq$ est manifestement impossible.

◇

¶C.3. La règle de COUPURE

La géométrie de cette règle dans les réseaux étant la même que celle du tenseur, la règle de coupure la plus libérale à ne donner que des réseaux corrects est donc la suivante:

$$\frac{\begin{array}{c} \vdots \pi_1 \\ \vdots \pi_2 \end{array} \quad \frac{\vdash K, A_1, A_2, \dots, A_n [i] \quad \vdash K^\perp, A_{n+1}, \dots, A_m [j]}{\vdash \bullet, A_1, A_2, \dots, A_n, A_{n+1}, \dots, A_m [\mathfrak{k}]} \text{COUPURE}}$$

¶C.4. Sémantique cohérente de la règle COUPURE

On définit la sémantique cohérente de cette règle comme d'habitude: $|\pi| = \{(a_1, \dots, a_n, a_{n+1}, \dots, a_m) / \exists z \in |K|(z, a_1, \dots, a_n) \in |\pi_1| \text{ et } (z, a_{n+1}, \dots, a_m) \in |\pi_2|\}$. Pour la même raison qu'habituellement s'il existe un tel z alors il est unique.

PROPOSITION 8.7. *Cette règle est interprétée par la sémantique cohérente.*

DÉMONSTRATION:

La preuve, on s'en doute, est grosso modo la même que pour la règle du TENSEUR, à ceci près que:

$$\begin{aligned} \bullet : \sim &\iff o : \wedge \\ \bullet : \wedge &\iff o : \sim \\ \bullet : = &\iff o : = \end{aligned}$$

On suppose que

$$(a_\bullet, a_1, a_2, \dots, a_n) \odot (a'_\bullet, a'_1, a'_2, \dots, a'_n) \left[\prod_i A_i \right]$$

et que

$$(a_o, a_{n+1}, \dots, a_m) \odot (a'_o, a'_{n+1}, \dots, a'_m) \left[\prod_j A_j \right]$$

avec $a_\bullet = a_o$ (on conserve les indices car $\circ : \wedge$ et $\bullet : \wedge$ n'ont pas la même signification) on va montrer que

$$(a_1, a_2, \dots, a_n, a_{n+1}, \dots, a_m) \odot (a'_1, a'_2, \dots, a'_n, a'_{n+1}, \dots, a'_m) \left[\prod_{\mathfrak{k}} A_k \right]$$

On écrira $l : \odot$ pour $a_l \odot a'_l$.

Les lettres i, i' désigneront des indices $\leq n$ ou l'indice $\bullet$ (correspondant à des éléments de $|i|$), et les lettres j, j' des indices $> n$ ou l'indice $\circ$ (correspondant à des éléments de $|j|$).

Si $(a_\bullet, a_1, a_2, \dots, a_n) = (a'_\bullet, a'_1, a'_2, \dots, a'_n)$ et $(a_o, a_{n+1}, \dots, a_m) = (a'_o, a'_{n+1}, \dots, a'_m)$ alors le résultat est évident.

Si $(a_\bullet, a_1, a_2, \dots, a_n) = (a'_\bullet, a'_1, a'_2, \dots, a'_n)$ et $(a_o, a_{n+1}, \dots, a_m) \wedge (a'_o, a'_{n+1}, \dots, a'_m)$, appelons j l'indice de j qui rendent les deux n-uplets strictement cohérents: $j : \wedge$ et pour tout j' de j au dessus de j dans j on a $j' :=$.

Cet indice est $j \neq \circ$, car $\bullet :=$ danc $\circ :=$, et on montre que ce même indice convient dans $\mathfrak{k}$. On a déjà $j : \wedge$. Soit $j' \neq \circ$ de j au dessus de j dans $\mathfrak{k}$; d'après la propriété (i') de $\mathfrak{k}$, cet indice j' est au dessus de j dans j , et donc $j :=$. Les $i \leq n$ sont tous $i :=$.

Le cas $(a_\bullet, a_1, a_2, \dots, a_n) \wedge (a'_\bullet, a'_1, a'_2, \dots, a'_n)$ et $(a_o, a_{n+1}, \dots, a_m) = (a'_o, a'_{n+1}, \dots, a'_m)$ est symétrique.

Si $(a_\bullet, a_1, a_2, \dots, a_n) \wedge (a'_\bullet, a'_1, a'_2, \dots, a'_n) \left[\prod_i A_i \right]$ et $(a_o, a_{n+1}, \dots, a_m) \wedge (a'_o, a'_{n+1}, \dots, a'_m) \left[\prod_j A_j \right]$

on a un indice i de $|i|$ qui rend les deux premiers n-uplets strictement cohérents et un indice j de j qui rend les deux seconds n-uplets cohérents. On va envisager trois cas (le cas $i = \bullet$ et $j = \circ$ ne se présentant pas), suivant que ces indices font ou non partie de $\bullet, \circ$.

Si $i = \bullet$ alors $\circ : \wedge$ et par suite $j \neq \circ$ et $\exists \circ > j$. Alors j convient. On a déjà $j : \wedge$. Soit $i \neq \bullet$ un point de $i - \bullet$ situé au dessus de j dans $\mathfrak{k}$. En raison de la propriété (iii) de $\mathfrak{k}$, comme $\circ$ n'est pas au dessus de j dans j ce point est au dessus de $\bullet$ dans i et donc $i :=$.

Si $i \neq \bullet$ et $j = \circ$. Sous-cas symétrique au précédent.

Si $i \neq \bullet$ et $j \neq \circ$, montrons que i ou j convient dans $\mathfrak{k}$. Si ce n'était pas le cas, il existerait j' de j au dessus de i dans $\mathfrak{k}$ tel que $j' \neq$, et un i' de i au dessus de j dans $\mathfrak{k}$ tel que $i' \neq$. le point i' ne serait donc pas au dessus de i dans $\mathfrak{k}$, et le point j' ne serait donc pas au dessus de j dans $\mathfrak{k}$. En vertu des propriétés (i) et (i') de $\mathfrak{k}$ le point i' ne serait pas au dessus de i dans i , et le point j' ne serait donc pas au dessus de j dans j ce qui contredirait la propriété (ii) de $\mathfrak{k}$. L'un des deux convient donc.

◇

§D. Quelques Propriétés de ce calcul

¶D.1. Toute règle précède à deux prémisses est simulable

Supposons qu'on ait une règle à deux prémisses pour le précède. Alors, par toute application de cette règle à une toute preuve π_1 de $A_\bullet, A_1, \dots, A_n[i]$ et à toute preuve π_2 de $A_o, A_{n+1}, \dots, A_m[j]$ on obtient une preuve π de $(A_\bullet, A_o), A_1, \dots, A_n, A_{n+1}, \dots, A_m[\ell]$. Appelons $\Pi_{1 \cup 2}$ le réseau constitué par la réunion des deux réseaux Π_1 et Π_2 associés aux deux preuves π_1 et π_2 . Notons que lorsque π_1 et π_2 varient les parcours de ce réseau décrivent tous les $u \cup v$, où $u \perp i$ et $v \perp j$. Considérons maintenant le module M de bord $A_\bullet, A_1, \dots, A_n, A_{n+1}, \dots, A_m[j]$ qui consiste en le lien $(A_\bullet, A_o)$ et l'ordre ℓ sur $(A_\bullet, A_o), A_1, \dots, A_n, A_{n+1}, \dots, A_m[\ell]$ et appelons $\mathcal{M}$ l'ensemble de ses parcours. Nous sommes dans les conditions du lemme 7.26. qui nous affirme alors que $\mathcal{M}^\perp = \{\ell / (A < B) \rightsquigarrow A \lesssim B\}^\perp$.

Dire que la règle est correcte, c'est dire que $\ell' = \ell / (A < B) \rightsquigarrow A \lesssim B$ est orthogonal à $u \cup v$ pour tout $u \perp i$ et tout $v \perp j$. C'est donc dire que ℓ' est une extension orthogonale de i, j . On peut donc dériver par la règle MÉLANGE le séquent $A_\bullet, A_o, A_1, \dots, A_n, A_{n+1}, \dots, A_m[\ell']$. On est alors en mesure d'appliquer notre règle à une prémisses PRÉCÈDE, puisque $A_\bullet \lesssim A_o$ dans $\ell' = \ell / (A < B) \rightsquigarrow A \lesssim B$ et le résultat est $(A_\bullet, A_o), A_1, \dots, A_n, A_{n+1}, \dots, A_m[\ell' / A \lesssim B \rightsquigarrow A < B]$. Il suffit alors de remarquer que $\ell = \left(\ell / (A < B) \rightsquigarrow A \lesssim B \right) / (A \lesssim B) \rightsquigarrow A < B$ pour se rendre compte qu'une telle règle n'augmente pas la classes des séquents prouvables.

¶D.2. Non commutation du TENSEUR et du MÉLANGE

Nous montrerons dans le paragraphe §C. du chapitre suivant qu'il est possible que la dernière règle soit TENSEUR (resp MÉLANGE) et l'avant dernière MÉLANGE (resp TENSEUR) sans que l'inverse soit possible, ce qui montre que ces deux règles ne commutent pas.

¶D.3. Elimination des coupures

Décrire les cas de base de l'élimination des coupures pour ce calcul n'est pas beaucoup plus difficile que pour les calculs plus simples; ce qui se passe sur les formules est identique; par contre, il faut s'assurer qu'on peut obtenir un ordre sur le séquent final de la preuve réduite restreint aux formules contienne celui de départ, ce qui, par une application de la règle ENTROPIE, permet de conclure. Pour pouvoir écrire le cas de base entre deux précède, il faut préalablement remarquer, que les quatre formules étaient dans des séquents disjoints, puisque qu'il faut que $A < B[i]$ pour pouvoir former la formule $(A < B)$, et que nos règles ne créent de l'ordre qu'entre des formules appartenant à des séquents différents (sinon on créerait à coup sur des circuits bigarrés dans les réseaux correspondants). On ne peut alors écrire ce cas qu'en supposant que les règles qui les ont amenés dans un même séquents précèdent immédiatement les règles PRÉCÈDE qui ont créé les deux formules coupées. Dans tous les cas le moyen le plus simple pour s'assurer que l'ordre final de la preuve réduite contient l'ordre de la preuve non réduite est de mimer ce qui se passe dans les réseaux, c'est-à-dire de s'assurer que le réseau formé est toujours correct, ce qui revient à montrer que l'orthogonal de l'ordre de la preuve réduite est inclus dans celui de la preuve non-réduite. On notera que c'est exactement ce qu'on a fait lorsqu'on a démontré la stabilité du critère de correction par élimination des coupures.

La preuve complète et directe de l'élimination des coupures est vraisemblablement affreuse. En effet, d'une manière générale, la complexité de cette démonstration réside surtout dans les commutations de règles, en particulier des règles structurelles (ici la règle de MÉLANGE), qui permettent de se ramener aux cas de base, et on a vu que les règles de ce calcul commutaient fort mal.

Chapitre 9

Séquentialisation

On ne fera pas de différence entre les liens tenseur et coupure qui ont et le même comportement combinatoire et des règles séquentielles similaires — comme il est habituel en ce domaine.

§A. Liens par et précède finaux

LEMME 9.1. *Un réseau ordonné Π de conclusions $(A_1 \wp A_2), A_3, \dots, A_n[i]$ est séquentialisable si et seulement si le réseau Π' de conclusion $A_1, A_2, \dots, A_n[i' = i / (A \wp B \rightsquigarrow A \sim B)]$ (obtenu par suppression de ce lien) est séquentialisable.*

DÉMONSTRATION: Il est clair que si Π' est séquentialisable, alors Π l'est aussi puisqu'on peut appliquer la règle du PAR ordonné à une preuve séquentielle Π' de Π' ($A \sim B$ dans i'), et qu'on obtient ainsi une preuve séquentielle dont le réseau correspondant est Π .

Soit maintenant π une séquentialisation de Π . Notons qu'une règle par correspondant au lien final $A_1 \wp A_2$ a été effectuée dans π , mais qu'il n'y a aucune raison qu'elle ait été effectuée en dernier. Néanmoins cette formule $A_1 \wp A_2$ n'est la prémisse principale d'aucune règle de π , et descend dans la preuve jusqu'au séquent conclusion. On va montrer que π se transforme canoniquement en une preuve π° dont le réseau associé est Π' . Remarquons que deux preuves dont chaque règle agit de la même manière sur les formules et dont l'ordre sur le séquent conclusion est le même ont le même réseau associé. La transformée π° de π sera une preuve telle qu'il soit possible de lui appliquer la règle PAR et d'obtenir ainsi une preuve dont le réseau soit Π , c'est à dire une preuve dont le réseau associé est Π' .

Pour obtenir π° on transforme chaque séquent de $\pi \vdash (A_1 \wp A_2), F_1, \dots, F_p[j]$ situé en dessous de la règle ayant créé $(A_1 \wp A_2)$ en un séquent $\vdash A_1, A_2, F_1, \dots, F_p[j / A_1 \wp A_2 \rightsquigarrow A_1 \sim A_2]$ ¹. Il faut encore vérifier qu'on obtient ainsi une preuve. L'action des règles sur les formules est évidemment correcte, et on vérifie aisément par induction sur la hauteur de la règle qui construit ce $A_1 \wp A_2$ dans la preuve que l'action des règles sur les ordres est aussi correcte.

Dans le cas d'une règle ENTROPIE, il suffit de remarquer que si $j' \subset j$ alors $j' / A_1 \wp A_2 \rightsquigarrow A_1 \sim A_2 \subset j / A_1 \wp A_2 \rightsquigarrow A_1 \sim A_2$.

¹ce genre de transformation n'est possible qu'en l'absence de règles de contraction, ce qui est le cas

Dans le cas d'une règle PAR ou PRÉCÈDE la remarque 1.6. (iv) du chapitre 1 montre que la règle est correcte. Dans le cas d'une règle d'ENTROPIE, c'est immédiat, il suffit d'affaiblir l'ordre de manière identique pour les deux points A et B .

Dans le cas d'une règle MÉLANGE, il suffit de remarquer que si $\mathfrak{k}$ est une extension orthogonale de $i \ j$ et que u est un point de i , alors $\mathfrak{k}/_{u \rightsquigarrow x \rightsquigarrow y}$ est une extension orthogonale de $i_{u \rightsquigarrow x \rightsquigarrow y} \ j$ (il y a même équivalence).

Enfin dans le cas d'une règle TENSEUR (dont $A_1 \wp A_2$ n'est évidemment pas une prémisse) il suffit de remarquer que si $\mathfrak{k}$ est une extension orthogonale de $i \ j$ au dessus de x et que u est un point de i différent de x , alors $\mathfrak{k}/_{u \rightsquigarrow x \rightsquigarrow y}$ est une extension orthogonale de $i_{u \rightsquigarrow x \rightsquigarrow y} \ j$ (il y a même équivalence). $\diamond$

LEMME 9.2. *Un réseau ordonné Π de conclusions $(A_1 \wp A_2), A_3, \dots, A_n[i]$ est séquentialisable si et seulement si le réseau Π' de conclusion $A_1, A_2, \dots, A_n[i' = i / (A \wp B \rightsquigarrow A \rightsquigarrow B)]$ (obtenu par suppression de ce lien) est séquentialisable.*

DÉMONSTRATION: Même démonstration que précédemment en remplaçant PAR et $\wp$ par $<$ et précède, et $\sim$ par $\lesssim$. $\diamond$

Remarquons que ces règles montrent que si une règle PAR (ou PRÉCÈDE) est effectuée avant une autre règle, alors il est possible de l'effectuer après.

On peut donc désormais supposer que les réseaux ordonnés considérés ont pour conclusions des formules tenseurs ou des atomes.

§B. Structure des réseaux sans par ni précède final

Les quelques définitions suivantes nous seront utiles pour décrire ces réseaux sans par ni précède final:

DÉFINITION 9.3. $\star$ On appellera ces réseaux sans par ni précède final des réseaux tendus.

$\star$ Remarquons que, par suppression des arcs, un réseau ordonné Π se transforme en un réseau Π' commutatif écrit avec la règle MÉLANGE (décrits et étudiés dans la partie B), que l'on appellera réseau commutatif associé.

$\star$ On appellera tenseurs scindants d'un réseau ordonné les tenseurs scindants du réseau commutatif associé (dont on a prouvé l'existence dans la partie 2).

$\star$ Un satellite d'un réseau ordonné est une composante connexe du réseau commutatif correspondant (tout point du réseau ordonné est dans un satellite, éventuellement réduit à ce

point). Notons qu'un arc d'un lien précède a toujours ses deux extrémités dans le même satellite.

* Un entrelacs d'un réseau ordonné est un ensemble de points d'un satellite défini comme une classe d'équivalence de la relation suivante: deux points X et Y d'un même satellite sont dits équivalents si et seulement si

▷ ils sont égaux ou

▷ aucun des deux n'est la conclusion d'un tenseur scindant, et, pour tout tenseur scindant du satellite, il sont situés d'un même coté de ce tenseur scindant dans le réseau commutatif associé.

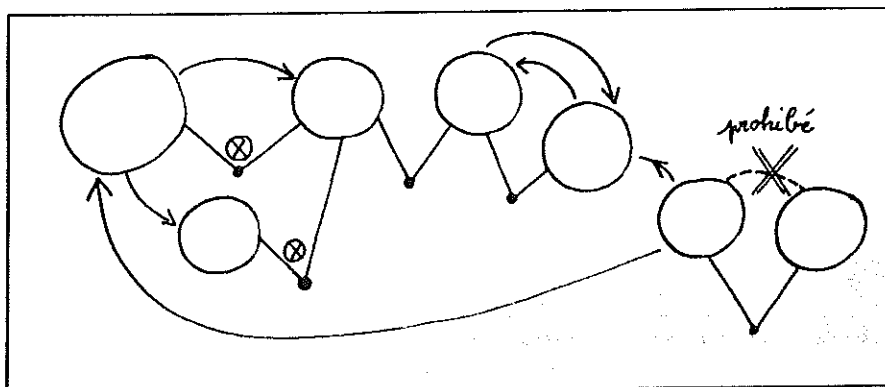
On notera qu'un tenseur scindant étant un isthme, cette définition a un sens.

DÉFINITION 9.4. On appelle agrégat orienté la donnée d'un agrégat et d'un ensemble d'arcs noirs reliant certains de ses sommets. On associe à un réseau ordonné un agrégat orienté sans circuit bigarré comme au chapitre 4, paragraphe §B.: les sommets de l'agrégat sont les arêtes axiomes du réseau tendu, et les couleurs propres ses liens tenseur. On a un arc noir de l'axiome a vers l'axiome b s'il existe un arc noir d'un point situé en dessous de a à un point situé en dessous de b .

On se convaincra qu'on obtient ainsi un agrégat orienté sans circuit bigarré, et qu'il n'y a pas lieu d'introduire de couleur pour les arcs, compte tenu qu'il sont orientés (!): si l'agrégat contenait un circuit bigarré, utilisant deux arcs de l'agrégat correspondant à un même arc du réseau, compte tenu de l'ordre de passage imposé (par le sens des arcs) ce circuit bigarré existerait aussi dans le réseau orienté.

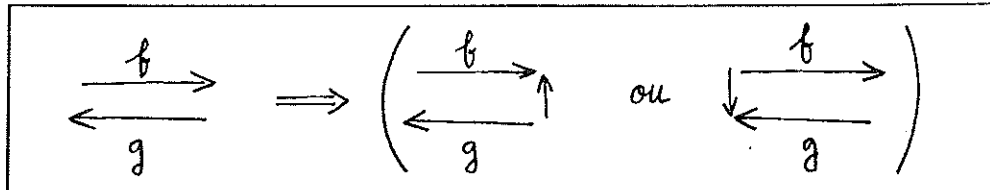
Les résultats des chapitres 2 et 4, montrent qu'un réseau commutatif consiste en une famille d'arbres (sans racines) dont les sommets sont de deux types: ce sont soit les conclusions des tenseur scindants, soit les entrelacs du réseau.

Un réseau ordonné est la superposition d'un réseau commutatif et d'arcs joignant ses sommets. De plus, le paragraphe §F. du chapitre 2, compte tenu de la fidélité de traduction de Π en $\Pi_{\otimes}$ (qui se vérifie comme au chapitre 4) on voit qu'il n'y a jamais d'arc entre deux feuilles situées à une distance 2 dans cette forêt. En effet on sait que tout point de la feuille est relié par un chemin bigarré bilatère à un tenseur scindant, et que la réunion de deux chemin bigarré ne traversant pas le même entrelacs est bigarré: l'existence d'un tel arc entraînerait donc l'existence d'un circuit bigarré. De même on peut affirmer qu'étant donné une feuille, il en existe une autre telle qu'il existe un chemin bigarré bilatère, entre tout point de la première et tout point de la seconde. Voici donc à quoi ressemble un réseau ordonné:



Nous avons utilisé dans les chapitres 1 et 8 les propriétés T et T_z , et nous allons en considérer une plus simple du même genre, qui sera plus maniable dans les graphes:

DÉFINITION 9.5. Deux arcs d'un graphe satisfont la condition $T^{\text{II}}(f, g)$ si et seulement si le graphe contient un arc de la source de l'un au but de l'autre:



On remarquera que nos graphes étant sans boucle, si ces deux arcs sont consécutifs, alors l'arc composé (obtenu par transitivité) fait nécessairement partie du graphe. Le lien entre cette condition et la transitivité explique le symbole choisi.

PROPOSITION 9.6. Supposons que le graphe soit un réseau ordonné et que ces deux arcs soient des arcs entre des conclusions; alors on peut supposer que si deux des quatre extrémités sont confondues, la condition est satisfaite.

DÉMONSTRATION: Notons que l'absence de boucle fait que ces quatre points sont au moins deux. Si ces quatre points ne sont en fait que deux alors ces deux arcs sont les mêmes, puisque le réseau contient au plus l'un des deux arcs (x, y) et (y, x) , et la condition est satisfaite. Si ces quatre points sont trois, et s'ils ont le même but ou la même source, alors la condition est satisfaite. Si ce sont des arcs consécutifs, comme le réseau est correct si et seulement si il l'est en ajoutant les arcs composés (puisque l'orthogonalité est stable par clôture transitive), on peut toujours supposer ces arcs présents, au quel cas la condition est satisfaite. $\diamond$

En vertu de cette remarque on supposera que les arcs composés sont toujours présents.

Donnons maintenant une définition qui va nous permettre de caractériser la dernière règle du calcul des séquents si elle est une règle de MÉLANGE :

DÉFINITION 9.7. Une frontière d'un réseau ordonné Π ayant au moins deux satellites est une partition en deux classes C_1, C_2 des satellites du réseau ordonné, ce qui fournit aussi une partition en deux classes des points du réseau.

Les arcs frontière sont les arcs qui la traversent, i.e. dont une extrémité est dans un satellite de C_1 et l'autre dans un satellite de C_2 ; ce sont donc des arcs entre conclusions. Deux arcs frontière sont dit contravariants si leurs extrémités finales sont de part et d'autre de la frontière.

Les extrémités des arcs frontière, appelés points frontière, se partagent en deux classes E_1 et E_2 de points conclusions. On notera qu'un chemin d'un point situé d'un côté de la frontière à un point situé de l'autre côté contient un nombre impair d'arcs frontière, et qu'un chemin d'un point à un point situé de l'autre côté de la frontière utilise un nombre pair d'arcs frontière.

La frontière est dite séparée si et seulement si la restriction de l'ordre du réseau aux points de la frontière satisfait la condition suivante, déjà apparue dans les chapitres 1 et 8:

(i) $T(f, E_1, E_2)$

Cela peut se formuler plus intelligiblement en termes d'arcs:

tout couple d'arcs frontière contravariants (f, g) satisfait $T^{\text{II}}(f, g)$.

On dit que la frontière est séparable s'il est possible d'ajouter des arcs à Π de sorte qu'on obtienne un graphe sans circuit bigarré et que la frontière soit séparée. Notons qu'il est alors possible de le faire de sorte que le graphe en question soit un réseau, puisqu'on n'a pas à ajouter d'arcs entre des points qui ne soient pas des conclusions du réseau pour satisfaire cette condition.

Donnons maintenant une définition du même genre qui va nous permettre de caractériser la dernière règle utilisée si elle est TENSEUR :

DÉFINITION 9.8. Soit Π un réseau tendu et $A \otimes B$ l'un de ses tenseurs scindants. Soit $\Pi - \{A \otimes B\}$ le réseau obtenu à partir de Π par suppression du point $A \otimes B$ et des arcs et arêtes qui lui sont incidents. Comme $A \otimes B$ est un tenseur scindant de Π les satellites de A et B dans $\Pi - \{A \otimes B\}$ sont distincts. Une frontière tendue d'un réseau ordonné Π (associée au tenseur scindant $A \otimes B$) est une partition en deux classes des satellites du réseau $\Pi - \{A \otimes B\}$ telle que le satellite de A soit dans une des deux classes et celui de B dans l'autre.

Un arc frontière est soit un arc qui la traverse, soit un arc incident à $A \otimes B$; un arc frontière est donc un arc entre deux conclusions.

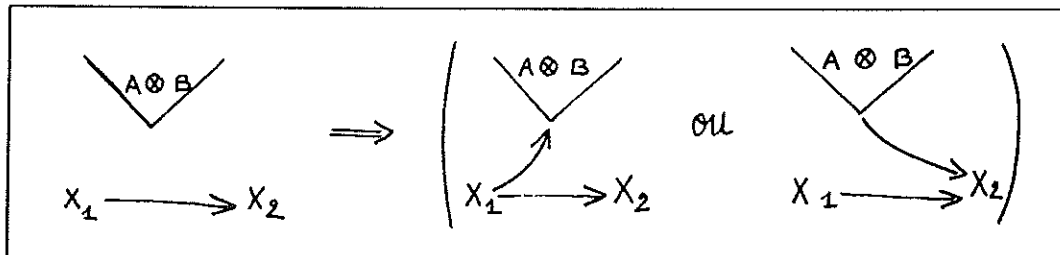
Les extrémités des arcs se partagent donc en trois classes $\mathcal{E}_1$ et $\mathcal{E}_2$ et $\{A \otimes B\}$, et ces points sont appelés points frontière. Deux arcs frontière sont dits contravariants si et seulement si: leurs sources ou leurs buts sont situées de part et d'autre de la frontière. On remarquera que, suivant cette définition les arcs $A \otimes B \rightarrow X_2 \in \mathcal{E}_2$ et $X'_1 \in \mathcal{E}_1 \rightarrow X'_2 \in \mathcal{E}_2$ sont contravariants, ce qui évitera la proliférations des clauses définissant une frontière tendue séparée.

Une frontière tendue est dite séparée si et seulement si la restriction de l'ordre du réseau aux points de la frontière satisfait la condition suivante déjà apparue dans les chapitre 1 et 8:

- (i) $\mathbb{T}_{A \otimes B}(f, \mathcal{E}_1, \mathcal{E}_2)$
- (ii) $\mathbb{T}(f, \mathcal{E}_1 \cup \{A \otimes B\}, \mathcal{E}_2)$
- (iii) $\mathbb{T}(f, \mathcal{E}_1, \mathcal{E}_2 \cup \{A \otimes B\})$

ce qui peut aussi se formuler en une liste plus intelligible de propriétés du graphe:

- (i) pour tout arc frontière $X_1 \in \mathcal{E}_1 \rightarrow X_2 \in \mathcal{E}_2$ le graphe contient l'un des deux arcs:
 $X_1 \in \mathcal{E}_1 \rightarrow A \otimes B$ ou $A \otimes B \rightarrow X_2 \in \mathcal{E}_2$;



et la condition symétrique: pour tout arc frontière $X_1 \in \mathcal{E}_1 \leftarrow X_2 \in \mathcal{E}_2$ le graphe contient l'un des deux arcs: $X_1 \in \mathcal{E}_1 \leftarrow A \otimes B$ ou $A \otimes B \leftarrow X_2 \in \mathcal{E}_2$.

- (ii) tout couple d'arcs frontière contravariants (f, g) satisfait $\mathbb{T}^{\text{II}}(f, g)$ (soulignons encore que l'une des quatre extrémités peut être $A \otimes B$).

On dit que la frontière tendue est séparable s'il est possible d'ajouter des arcs à Π de sorte qu'on obtienne un réseau et que la frontière soit séparée.

Maintenant mentionnons deux lemmes qui montrent que les définitions précédentes de frontière séparable et de frontière tendue séparable fournissent l'analogue des habituels tenseurs scindants i.e. caractérisent les dernières règles appliquées.

Le signe Π° désigne désormais un réseau tels que tous les réseaux de taille inférieure à la sienne sont séquentialisables.

LEMME 9.9. *Soit $\mathcal{F} = (C_1, C_2)$ une frontière d'un réseau Π° . Alors $\mathcal{F}$ séparable équivaut à: il existe une preuve séquentielle de réseau associé Π° dont la dernière règle est MÉLANGE suivie éventuellement de ENTROPIE appliquée à une séquentialisation de C_1 dans un ordre i_1 et à une séquentialisation de C_2 dans un ordre i_2 .*

DÉMONSTRATION: L'implication réciproque est triviale. Pour l'implication directe il suffit de remarquer que pour que l'ordre satisfasse (i) il n'est pas nécessaire d'ajouter des arcs ailleurs qu'entre deux conclusions. $\diamond$

LEMME 9.10. *Soit $\mathcal{F} = (A \otimes B, C_1, C_2)$ une frontière tendue d'un réseau Π° . Alors $\mathcal{F}$ séparable équivaut à: il existe une preuve séquentielle de réseau associé Π° dont la dernière règle est TENSEUR suivie éventuellement de ENTROPIE appliquée à une séquentialisation de C_1 dans un ordre i_1 et à une séquentialisation de C_2 dans un ordre i_2 .*

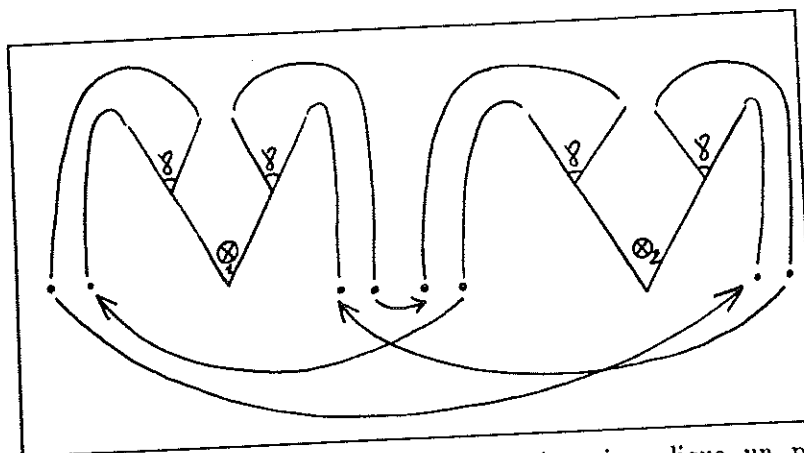
DÉMONSTRATION: Aussi évidente que la précédente. $\diamond$

§C. Monstres Choisis

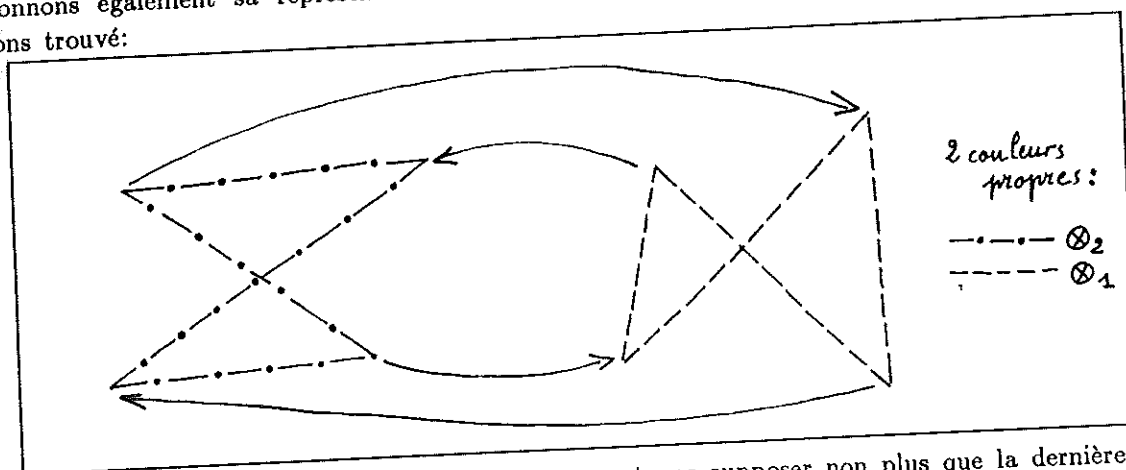
Avant d'en venir à la séquentialisation proprement dite, nous souhaitons présenter ici quelques uns des exemples qui nous ont permis d'élaborer le calcul des séquents, et de trouver la méthode de séquentialisation, ainsi que de démontrer que les règles TENSEUR et MÉLANGE ne commutent pas.

Nous avons premièrement considéré règle de mélange plus stricte, où l'ordre sur le séquent conclusion était $i < j$, mais on s'est aperçu qu'elle était trop stricte pour permettre la séquentialisation, même avec la règle la plus libérale possible pour le tenseur (celle donnée au chapitre précédent). Le réseau qui suit a les propriétés suivantes:

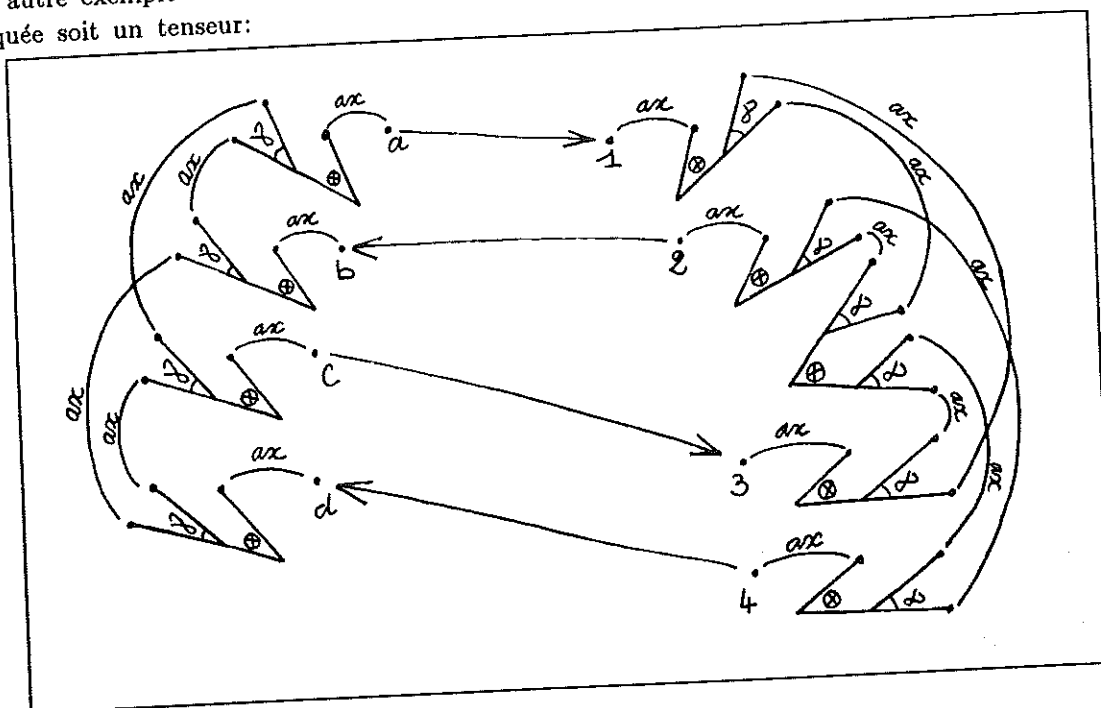
- * aucun de ses deux tenseurs scindants n'admet de frontière tendue séparable, et sa dernière règle n'est donc pas un TENSEUR ; ce qui montre qu'on ne peut se restreindre à chercher la dernière règle sous la forme d'un TENSEUR
- * sa seule frontière est traversée dans les deux sens, ce qui montre qu'on ne peut pas se contenter d'une règle MÉLANGE où l'ordre sur le séquent conclusion est $i < j$; par contre cette frontière est séparable, et une fois le réseau ainsi séparé, chacun des deux tenseurs admet une frontière tendue séparable dans sa partie.
- * cet exemple est minimal pour ces propriétés



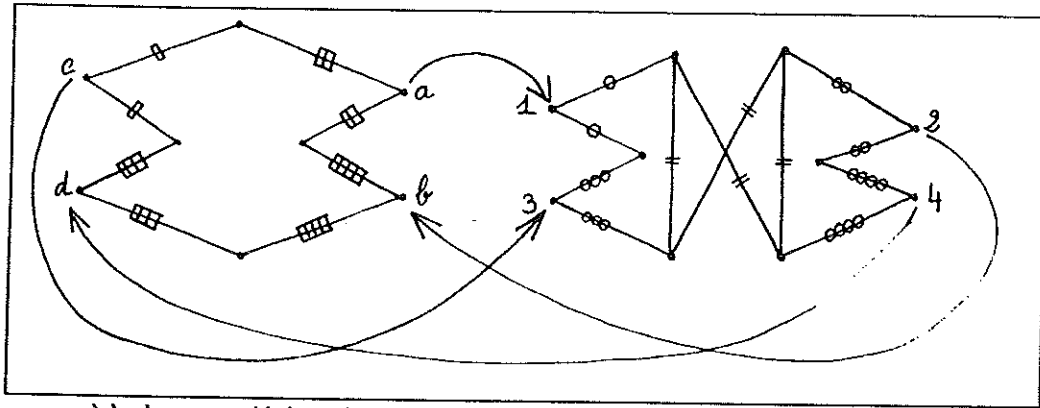
Donnons également sa représentation en agrégat orienté, qui explique un peu comment nous l'avons trouvé:



Un autre exemple vient alors montrer qu'on ne peut pas supposer non plus que la dernière règle appliquée soit un tenseur:



donnons aussi sa représentation en agrégat orienté:



Ce réseau possède les propriétés suivantes:

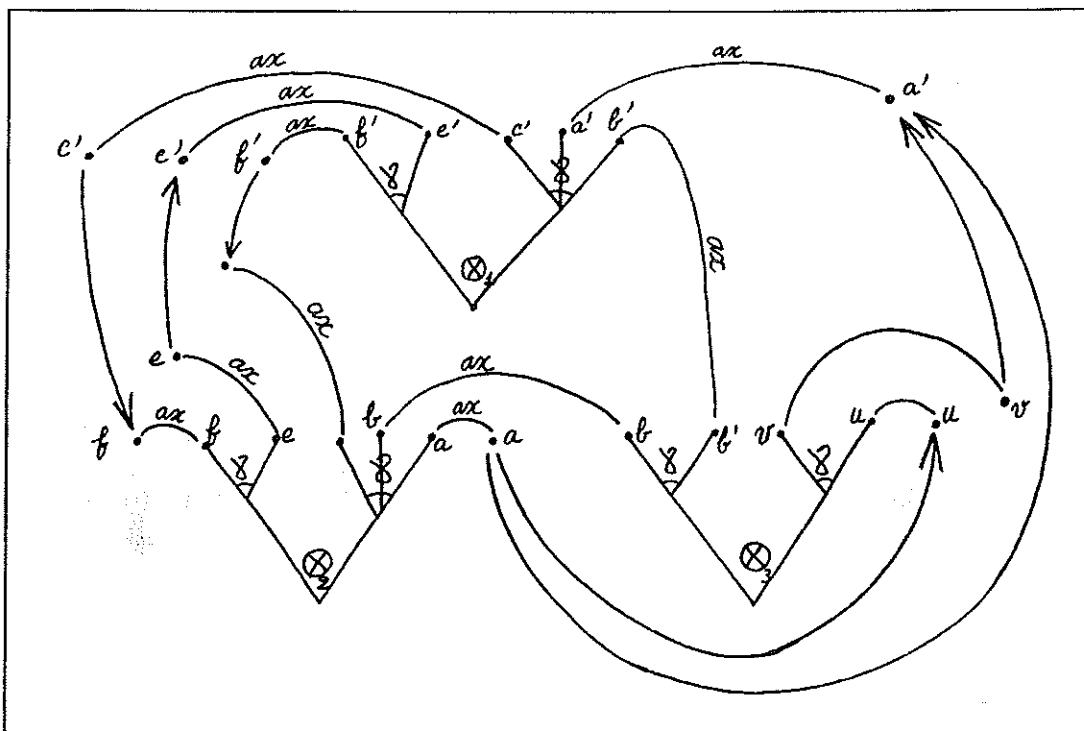
- * L'unique frontière possible n'est pas séparable, ce qui montre qu'on ne peut supposer que la dernière règle appliquée soit une règle de MÉLANGE .
- * Aucun arc frontière ne peut être transformé en tenseur sans créer de circuit bigarré. S'il n'existait pas de tels réseaux, on eut alors pu supposer que la dernière appliquée est toujours un TENSEUR , en changeant éventuellement des arcs en tenseurs. (C'est ce qui explique qu'on l'ait préféré au contre-exemple minimal à satisfaire le premier point)

Nous avons cru un certain temps que les réseaux ordonnés étaient séquentialisables avec une règle du tenseur plus stricte que celle présentée dans le chapitre précédent, disons TENSEUR⁻ où l'on demande seulement que dans l'ordre sur le séquent conclusion aucune formule du second séquent prémisses ne soit inférieure à une formule du premier séquent prémisses — c'est bien évidemment un cas particulier de notre règle.

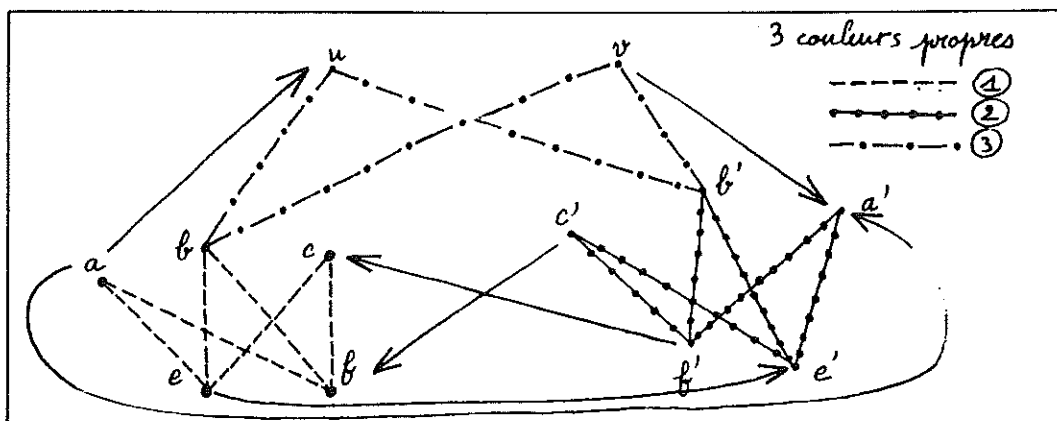
La dernière règle du calcul des séquents, si elle est une telle règle TENSEUR⁻, correspond à une frontière qui n'est traversée que dans un sens. Notons de plus que, si un réseau tendu ordonné n'a qu'un seul satellite, alors la dernière règle est nécessairement un TENSEUR , et, que pour un lien tenseur scindant donné, il n'y a qu'une seule frontière tendue possible. Mais nous avons trouvé alors un réseau avec les propriétés suivantes:

- * ce réseau n'a qu'un seul satellite
- * sa dernière règle ne peut pas être TENSEUR⁻ puisque toute frontière tendue est traversée dans les deux sens *et il est minimal pour cette propriété*
- * un seul de ses trois tenseurs scindants peut être la dernière règle (celui du milieu), ce qui montre nos règles commutent mal, et qu'un tenseur scindant peut ne jamais correspondre à la dernière règle appliquée.

que voici:



Donnons également sa représentation en agrégat orienté :



Nous pensons que sans cette notion d'agrégat il eut été inespéré de trouver un tel contre-exemple dont on a pourtant démontré qu'il était minimal.

Nous n'avons pas eu le courage d'écrire les preuves séquentielles correspondantes

§D. Séquentialisation

Compte tenu des monstres ci-dessus présentés, nous pouvons entamer la démonstration de la séquentialisation en gardant présent à l'esprit qu'on ne peut se restreindre à chercher exclusivement une frontière séparable ou à chercher exclusivement une frontière tendue séparable. Il faut donc

chercher simultanément l'une ou l'autre. C'est là une différence essentielle avec le calcul de la seconde partie.

DÉFINITION 9.11. Une *saturation* d'un réseau ordonné Π est un graphe simple sans circuit bigarré obtenu par addition d'arcs noirs à Π tel que, étant donnés deux points A et B du graphe, si ce graphe ne contient pas l'arc $B \rightarrow A$ alors il y a un chemin bigarré de A à B . Dans un tel graphe, l'addition d'un arc crée toujours un circuit bigarré, d'où le nom "saturation".

PROPOSITION 9.12. Tout réseau ordonné admet au moins une saturation.

DÉMONSTRATION: On le batit de proche en proche par addition d'arcs noirs, sans difficulté aucune: on construit une suite de graphe de ce type dont le premier est $\Pi_0 = \Pi$; si Π_n contient deux sommets A et B tel que Π_n ne contienne pas de chemin bigarré de A vers B , ni l'arc $B \rightarrow A$, on prend Π_{n+1} égal à Π_n augmenté de l'arc noir $B \rightarrow A$ et Π_{n+1} ne contient pas plus de circuit bigarré que Π_n , c'est-à-dire aucun. Notons que par la suite nous n'aurons pas à reconsidérer ce couple. Une telle suite s'arrête donc, s'il y a s sommets, en moins de $s(s-1)/2$ fois. $\diamond$

Comme il est impossible d'ajouter un arc à une saturation, on a:

REMARQUE 9.13. Dans une saturation d'un réseau, les notions de frontière (resp frontière tendue) séparable et de frontière (resp frontière tendue) séparée coïncident.

LEMME 9.14. Il existe une frontière séparable dans Π équivaut à il existe une saturation Π^+ de Π dans la quelle il existe une frontière séparée.

DÉMONSTRATION:

* S'il existe une frontière séparée dans une saturation de Π , en omettant les arcs non nécessaires au fait que cette frontière soit séparée, on obtient une frontière saturée de Π , obtenue par adjonction d'arcs noirs à Π , ce qui montre que la frontière est séparable.

* S'il existe une frontière séparable C_1, C_2 dans Π , on peut par définition adjoindre des arcs à Π de sorte à ce qu'elle soit séparée. Appelons Π^* le réseau obtenu et montrons qu'il est possible de le saturer de sorte que la frontière reste séparée.

Si, pour le saturer, on ajoute à Π_n^* un arc qui n'est pas un arc frontière, la dite frontière reste saturée et on poursuit la saturation.

Si pour le saturer on doit ajouter un arc frontière, c'est qu'il existe un point A de C_1 et un point B de C_2 tels qu'il n'existe pas de chemin bigarré de B à A dans Π_n^* , et Π_{n+1}^* est obtenu par addition de l'arc $f = A \rightarrow B$. Les seuls couples d'arcs frontière qui empêcheraient éventuellement la frontière d'être séparée sont les couples f, g où g est un arc frontière contravariant avec f , $g = A' \leftarrow B'$. S'il n'est pas possible d'ajouter à Π_{n+1}^* ni l'arc $B' \rightarrow B$ ni l'arc $A \rightarrow A'$ c'est que Π_{n+1}^* contient un chemin bigarré u de B' à B et un chemin bigarré v de A à A' , et prenons les de longueur minimale. Comme il n'y a pas de circuit bigarré, ces deux chemins utilisent donc des arêtes de même couleur, et par suite, l'un au moins des deux traverse la frontière, disons u . Soient f' et g' deux arcs contravariants de la frontière consecutifs sur u ; comme ces arcs sont des arcs de Π_n^* , dans le quel cette frontière est séparée, on a $T^{\text{ll}}(f, g)$ et on voit que u ne peut être de longueur minimale. L'un de ces deux chemins n'existe donc pas et on peut ajouter

l'un des deux arcs à Π_{n+1} , et on fait cela pour tout g tel que f, g ne satisfasse pas la condition. Ensuite on poursuit la saturation, avec une frontière à nouveau séparée.

◇

LEMME 9.15. *Il existe une frontière tendue séparable dans Π équivaut à il existe une saturation Π^+ de Π dans la quelle il existe une frontière tendue séparée.*

DÉMONSTRATION: Preuve similaire à la précédente.

◇

On en déduit aisément le

LEMME 9.16. *Un réseau contient une frontière séparable ou une frontière tendue séparable si et seulement s'il existe une saturation de ce réseau qui contient une frontière séparée ou une frontière tendue séparée.*

Et maintenant nous donnons une caractérisation des frontières non-séparées dans une saturation:

THÉORÈME 9.17. *Une frontière $\mathcal{F}$ d'une saturation n'est pas séparée si et seulement s'il existe deux arcs frontière contravariant f et f' et un circuit $X(f \in \mathcal{F})X'pY'(f' \in \mathcal{F})X'qX$ où le chemin p est bigarré et n'utilise pas d'arc de $\mathcal{F}$, et où le chemin q est bigarré.*

DÉMONSTRATION: Montrons pour commencer que s'il existe un tel circuit la frontière n'est pas séparée: il suffit de considérer le couple f, f' pour s'en convaincre: ce couple ne peut satisfaire $T^{\downarrow}(f, f')$.

Réciproquement si la frontière n'est pas séparée, alors il existe un couple d'arcs traversant la frontière $A_1 - f \rightarrow A'_1$ et $A_2 \leftarrow f' - A'_1$ tels qu'il existe un chemin bigarré p de A'_1 à A'_2 et un chemin bigarré p' de A_2 à A_1 . Ce circuit traverse nécessairement la frontière d'autre fois sinon p et p' n'aurait pas de couleur commune, et la saturation contiendrait un circuit bigarré.

Appellons distance entre f et f' sur un circuit $fuf'vf$ le minimum de n_u et n_v où n_u (resp n_v) désigne le nombre d'arcs de la frontière contenus dans u (resp v) — ce nombre est évidemment pair.

Prenons maintenant un circuit traversant contenant f et f' tel que leur distance sur ce circuit soit minimale. Considérons maintenant deux arcs de la frontière consécutifs sur un ce circuit, du côté où est atteinte cette distance (s'il n'y en a pas, la distance est 0 et le problème est résolu) disons $X - g \rightarrow X'$ et $Y \leftarrow g' - Y'$. On a donc un chemin bigarré d de Y à X ne traversant pas la frontière. Si l'agrégat contient l'arc $Y' - u \rightarrow X'$, alors on a trouvé un circuit sur lequel la distance entre f et f' est plus petite que sur le circuit de départ, ce qui est impossible. Par définition d'une saturation, il existe un chemin bigarré e de X' à Y' . Les arcs g et g' sont donc sur un circuit $X(g)X'eY'(g')X'dX$ où e et d sont bigarrés et où d ne traverse pas la frontière.

◇

Pour donner une caractérisation des frontières tendues non-séparées dans une saturation, nous utiliserons le petit lemme suivant sur les frontières séparées:

LEMME 9.18. *Soit $\mathcal{F} = C_1, C_2$ une frontière séparée.*

- (i) *soit X et Y deux points de C_1 (ou de C_2) tels qu'il existe un chemin bigarré de X à Y ; alors il existe un chemin bigarré de X à Y n'utilisant pas d'arcs frontière.*

- (ii) soient X un point de C_1 et Y un point de C_2 (ou X un point de C_2 et Y un point de C_1) tels qu'il existe un chemin bigarré de X à Y ; alors il existe un chemin bigarré de X à Y utilisant un et un seul arc frontière.

DÉMONSTRATION:

- (i) Considérons un chemin bigarré de X à Y utilisant un nombre minimal d'arcs frontière. On sait que ce nombre est pair. Si c'est 0, le lemme est établi. Sinon, soit deux arcs frontière f, g consécutifs sur ce chemin; ils sont contravariant et satisfont donc $T^{\text{II}}(f, g)$. Si ces arcs sont $f = A_1 \rightarrow A_2$ qui précède $g = A'_1 \leftarrow A'_2$, comme la frontière est séparée le graphe contient l'un des deux arcs $A_1 \rightarrow A'_1$ ou $A_2 \leftarrow A'_2$, et comme ce ne peut être $A_2 \leftarrow A'_2$ (sinon il y aurait un circuit bigarré), on a l'arc $A_1 \rightarrow A'_1$, ce qui contredit la minimalité du chemin bigarré considéré.
- (ii) Argument similaire.

◇

THÉORÈME 9.19. Une frontière tendue $\mathcal{F}$ d'une saturation n'est pas séparée si et seulement s'il existe au moins l'une des deux configurations suivantes:

- (i) deux arcs frontière contravariants f et f' et un circuit $AdX(f)X'eY'(f')YgA$ (ou la même chose avec B) tel que f et f' sont les deux seuls arcs frontière de ce circuit, les chemins $AdX(f)X'eY'$ et $X'eY'(f')YgA$ sont bigarrés.
- (ii) deux arcs frontière contravariants f et f' (non incident au tenseur) et un circuit $X(f \in \mathcal{F})X'pY'(f' \in \mathcal{F})X'qX$ où le chemin p est bigarré et n'utilise pas d'arc de $\mathcal{F}$, et où le chemin q est bigarré (comme précédemment)

DÉMONSTRATION:

- ★ Il est clair que s'il on a l'une des deux configurations, alors la frontière tendue n'est pas séparée, la première configuration contredisant le point (i) de la définition de frontière tendue séparée, et le second contredisant le point (ii) de cette définition.
- ★ On suppose que la frontière n'est pas saturée, et qu'on n'a pas (ii); on peut donc appliquer T^{II} aux couples d'arcs frontière contravariants non-incidents à $A \otimes B$, et en particulier supposer qu'un chemin traversant en utilise au plus un. On suppose qu'on a $f = X \rightarrow Y$ sans avoir ni $X \rightarrow A \otimes B$ ni $A \otimes B \rightarrow Y$; puisque le graphe est saturé, il existe un chemin bigarré d de Y à $A \otimes B$ et un chemin bigarré e de $A \otimes B$ à X . Si ni e ni chd n'empruntaient les arêtes du tenseur t , on aurait un circuit bigarré: ce cas est donc exclus. Si e et chd utilisent t , la propriété précédente montre qu'on peut supposer que e et d n'utilisent qu'un seul arc frontière; la condition T^{II} montre qu'alors qu'on a la configuration (ii). Si l'un des deux n'utilise pas t , alors l'un traverse une seule fois et l'autre aucune et on a aussi une configuration (ii).

◇

On remarque que ces théorèmes entraînent l'existence de chemins bigarrés quasi compatibles, et ce pour toute frontière non-séparée, et pour toute frontière tendue non-séparée.

Pour démontrer la séquentialisation, nous utiliserons un lemme, que nous n'avons pas encore fini de démontrer mais qui est très raisonnable au vu des deux théorèmes précédents; nous envisagerons cependant ce que signifierait l'échec de la séquentialisation, car en ce domaine l'intuition peut facilement être prise en défaut:

LEMME 9.20. *(non encore établi) Si une saturation ne contient ni frontière séparée ni frontière tendue séparée, alors elle contient un circuit bigarré.*

La raison pour la quelle nous croyons cela est que les lemmes précédents montre qu'on a alors un grand nombre de chemins dont la réunion est bigarrée, et nos travaux en cours nous encouragent dans cette croyance .

Moyennant ce lemme, on peut démontrer la séquentialisation:

THÉORÈME 9.21. *(modulo le lemme précédent) Tout réseau ordonné est séquentialisable.*

DÉMONSTRATION: Considérons un réseau Π° non-séquentialisable de taille minimale — i.e. tous les réseaux ayant moins de liens sont séquentialisables. En vertu du premier paragraphe de ce chapitre, Π° ne contient ni par ni précède final. S'il n'est pas séquentialisable, c'est donc qu'il ne contient aucune frontière tendue séparable, et aucune frontière séparable. Considérons alors une saturation de Π , et on a vu précédemment qu'il en existe. Alors cette saturation ne contiendrait aucune frontière séparée et aucune frontière tendue séparée, et notre pseudo-lemme précédent montre que c'est impossible. $\diamond$

Qu'advierait-il si l'on découvrait un réseau ordonné non séquentialisable?

A la différence d'autres travaux où l'on n'a pas de preuve de la séquentialisation (connecteurs généralisés de [DR90]), nous proposons ici un calcul des séquents, qui, de plus, est l'unique candidat possible (en vertu du chapitre 8). Cela signifierait qu'il n'y a pas de calcul des séquents (raisonnable) correspondant à ce calcul en réseaux, qui pourtant étend le plus conservativement possible le calcul usuel — notamment pour la sémantique cohérente.

En fait, si ce n'était de l'énergie dépensée pour essayer de démontrer ce résultat, on pourrait même se réjouir de découvrir ainsi de nouvelles directions.

En guise de conclusion

Comme on l'aura remarqué, notre étude du calcul ordonné est loin d'être exhaustive.

Une question naturelle pour un système logique est celle de sa complétude: il semble que le moyen le plus raisonnable de mener cette question à bien soit d'étendre la sémantique des phases, en définissant une opération interprétant les connecteurs généralisés et le *précède*.

La décidabilité de la prouvabilité est aussi une question naturelle: nous pensons, au vu du calcul en réseaux qu'elle ressemble à celle du calcul purement multiplicatif habituel.

Mentionnons encore que le développement de l'interprétation informatique, en se référant par exemple aux travaux de [MTV90] nous semble intéressant.

Parmi les extensions de ce calcul, nous avons déjà étudié celle à un calcul avec les modalités "?" et "!" et les quantificateurs, pour obtenir un calcul d'un pouvoir d'expression raisonnable, sans boîtes affaiblissement comme la présence de la règle MÉLANGE nous le permet; le calcul ainsi obtenu est donc confluent. Les règles doivent opérer ainsi sur l'ordre: lors d'une règle correspondant à la formation d'une boîte, l'ordre sur les conclusions est le même qu'avant la formation de cette boîte; la règle de contraction ne doit contracter que des formules équivalentes dans l'ordre (c'est grosso modo un *par*). Lors de l'élimination des coupures, chaque formule du contexte d'une boîte ! dupliquée, est dupliquée en deux formules équivalentes dans l'ordre. On peut alors démontrer la normalisation forte de ce calcul, comme dans [Gir87a] complété par la propriété de striction de [Dan90]; et la confluence résulte de l'absence des boîtes affaiblissement.

Une autre généralisation possible serait de travailler avec des préordres plutôt qu'avec des ordres, ce qui revient à autoriser des circuits bigarrés internes à *i*; cette généralisation semble à première vue assez compliquée à réaliser.

Une extension intéressante serait de formuler un calcul des séquents contenant et des formules à gauche et des formules à droite. Nous pensons qu'il s'agit là d'une question très difficile, vu que le dual d'un espace cohérent défini par un ordre n'est pas, en général, un espace cohérent défini par un ordre, sauf dans le cas où l'ordre est total, ce qui extrêmement pathologique: c'est un cas où toutes les preuves des différentes formules sont complètement déconnectées. Notons tout de même que dans le cas d'un ordre contractile, on sait définir ce dual, puisque la formule s'écrit avec les connecteurs *par* et *précède* dont on connaît le dual.

Signalons enfin qu'il nous semblerait judicieux de remplacer les réseaux, d'une définition propre à endormir tout non-logicien, par les agrégats et d'ainsi pouvoir appliquer des techniques (sinon des résultats) combinatoires standard. Nous avons été surpris de l'absence de connection entre l'étude de la structure des preuves et les mathématiques combinatoires, alors que les preuves sont évidemment des structures finies, généralement des arbres, et dans le cas des réseaux de la logique linéaire des graphes colorés — que nous préférons aux positions d'interrupteurs, même s'il s'agit bien des mêmes objets....

1. The first conclusion is that the system is not stable.

2. The second conclusion is that the system is not stable.

3. The third conclusion is that the system is not stable.

4. The fourth conclusion is that the system is not stable.

5. The fifth conclusion is that the system is not stable.

6. The sixth conclusion is that the system is not stable.

7. The seventh conclusion is that the system is not stable.

8. The eighth conclusion is that the system is not stable.

9. The ninth conclusion is that the system is not stable.

10. The tenth conclusion is that the system is not stable.

11. The eleventh conclusion is that the system is not stable.

12. The twelfth conclusion is that the system is not stable.

13. The thirteenth conclusion is that the system is not stable.

14. The fourteenth conclusion is that the system is not stable.

15. The fifteenth conclusion is that the system is not stable.

16. The sixteenth conclusion is that the system is not stable.

17. The seventeenth conclusion is that the system is not stable.

18. The eighteenth conclusion is that the system is not stable.

19. The nineteenth conclusion is that the system is not stable.

20. The twentieth conclusion is that the system is not stable.

21. The twenty-first conclusion is that the system is not stable.

22. The twenty-second conclusion is that the system is not stable.

23. The twenty-third conclusion is that the system is not stable.

24. The twenty-fourth conclusion is that the system is not stable.

25. The twenty-fifth conclusion is that the system is not stable.

26. The twenty-sixth conclusion is that the system is not stable.

27. The twenty-seventh conclusion is that the system is not stable.

28. The twenty-eighth conclusion is that the system is not stable.

29. The twenty-ninth conclusion is that the system is not stable.

30. The thirtieth conclusion is that the system is not stable.

31. The thirty-first conclusion is that the system is not stable.

32. The thirty-second conclusion is that the system is not stable.

33. The thirty-third conclusion is that the system is not stable.

34. The thirty-fourth conclusion is that the system is not stable.

35. The thirty-fifth conclusion is that the system is not stable.

Bibliographie

- [Abr91] Samson Abramsky. Computational interpretation of linear logic. Technical report, Imperial College, London, 1991. to appear in TCS.
- [Ber73] Claude Berge. *Graphs and Hypergraphs*. North-Holland, 1973.
- [Ber79] Claude Berge. *Graphes*. Gauthier-Villars, 1979.
- [Bol79] Béla Bollobás. *Graph Theory*, volume 63 of *Graduate Texts in Mathematics*. Springer Verlag, 1979.
- [Dan90] Vincent Danos. *La logique linéaire appliquée à l'étude de divers processus de normalisation et principalement du λ -calcul*. PhD thesis, Université Paris7, Mathématiques, 1990.
- [DR90] Vincent Danos and Laurent Regnier. The structure of multiplicatives. *Archives for mathematical logic*, 28:181–203, 1990.
- [FR90] Arnaud Fleury and Christian Retoré. The mix rule. Prépublications de l'équipe de Logique 11, U.F.R. de Mathématiques, Université Paris 7, 1990. to appear in Mathematical Structures in Computer Science.
- [Gir87a] Jean-Yves Girard. Linear logic. *Theoretical Computer Science*, 50, 1987.
- [Gir87b] Jean-Yves Girard. Multiplicatives. In *Proceedings of the Congress of Logic and Computer Science held in Torino, Octobre 1986*, 1987.
- [Gir90] Jean-Yves Girard. La logique linéaire. *Pour La Science, Edition Française de "Scientific American"*, (150), April 1990. English translation to appear in Scientific American.
- [GLT88] Jean-Yves Girard, Yves Lafont, and Paul Taylor. *Proofs and Types*. Number 7 in Cambridge Tracts in Theoretical Computer Science. Cambridge University Press, 1988.
- [GW75] Graver and Watkins. *Combinatorics with Emphasis on the Theory of Graphs*, volume 54 of *Graduate Texts in Mathematics*. Springer Verlag, 1975.
- [How80] William A. Howard. The formulae-as-types notion of construction. In J. Hindley and J. Seldin, editors, *To H.B. Curry: Essays on Combinatory Logic, λ -calculus and Formalism*, pages 479–490. Academic Press, 1980.
- [Kri90] Jean-Louis Krivine. *Lambda Calcul — Types et Modèles*. Etudes et Recherches en Informatique. Masson, Paris, 1990.
- [LW92] Patrick Lincoln and T. Winkler. Constant-Only Multiplicative Linear Logic is NP-complete. manuscript, 1992.

- [Mét92] François Métayer. Homologie des réseaux. Prépublications de l'équipe de logique, U.F.R. de Mathématiques, Université Paris 7, 1992.
- [MTV90] Marcel Masseron, Christophe Tollu, and Jacqueline Vauzeilles. Plan generation and linear logic. In *FST-TCS 10*, Lecture Notes in Computer Science, pages 63–75. Springer-Verlag, 1990.
- [Par88] Michel Parigot. Preuves et programmes: les mathématiques comme langage de programmation. *Le courrier du C.N.R.S.*, 1988. Images des Mathématiques, supplément au numéro 69.
- [Tro92] Anne Sjerp Troelstra. *Lectures on Linear Logic*, volume 29 of *CLSI Lecture Notes*. Stanford University Press, 1992.
- [vdW89] Jacques van de Wiele. De l'utilisation des jeux en logique linéaire. notes circulant dans l'Equipe de Logique, Université Paris 7, 1989.
- [Yet90] David N. Yetter. Quantales and (non-commutative) linear logic. *Journal of Symbolic Logic*, 55(1), March 1990.

Index des notations

Connecteurs de la logique linéaire		
Girard (et nous)	Troelstra	
$\wp$	$+$	par ("ou" multiplicatif)
$\perp$	0	faux (neutre du par)
$\otimes$	$*$	tenseur ("et" multiplicatif)
1	1	un (neutre du tenseur)
$A^\perp$	$\sim A$	A orthogonal (négation linéaire)

Petites majuscules	TENSEUR	règles du calcul des séquents
Caractères dactylographiés	tenseur	connecteurs et liens correspondants
Majuscules romaines	A,..U,..X..	formules
		sommets d'un graphe
Minuscules romaines	a,b,..	points d'une relation
		arêtes d'un graphes
		parfois, sommets d'un graphe
Majuscules calligraphiées	$\mathcal{H}, \mathcal{G} \dots$	graphes, agrégats
	$\mathcal{N}$	la couleur noire
	C,E..	cotés d'une frontière
	$\mathcal{F}$	parcours du module F^∇ associé à F
Minuscules sans sérif	c d	chemins d'un graphes
Minuscules gothiques	u ,v ,w ...	relations
	i ,j ,k ,l ...	ordres (partiels)
Minucules grecques	$\alpha, \beta, \dots$	couleurs propres
	π	preuve séquentielle
Majuscules grecques	Π	réseaux, ordonnés ou commutatifs
Majuscules doubles	P,C,T	conditions

$\mathcal{I}$	17 relation identique
$ u $	17 domaine de la relation u
$u _E$	17 restriction de la relation u à l'ensemble E
$\bar{u}$	17 clôture transitive et réflexive de la relation u
$x < y[i]$	17 x strictement inférieur à y dans i
$x \leq y[i]$	17 x inférieur à y dans i
$x \sim y[i]$	17 points équivalents dans i
$x \lesssim y[i]$	18 points équivalents-inférieurs dans i
$i/u \rightsquigarrow x \rightsquigarrow y$	18 contraction d'un ordre suivant deux points équivalents
$i/x \lesssim y \rightsquigarrow u$	19 contraction d'un ordre suivant deux points équivalents-inférieurs
$i/x \rightsquigarrow y \rightsquigarrow u$	21 expansion d'un ordre suivant deux points équivalents
$i/u \rightsquigarrow x < y$	21 expansion d'un ordre suivant deux points équivalents-inférieurs
$\mathbb{P}$	22 condition pour qu'un ordre soit contractile
$u \perp v$	27 u et v sont deux relations faiblement orthogonales
$u \perp v$	28 u et v sont deux relations orthogonales
$T(I, E, E')$	29 condition pour qu'un ordre soit une extension orthogonale de deux ordres disjoints
$T_x(I, E, E')$	30 condition pour qu'un ordre soit une extension orthogonale de deux ordres ayant un point commun
C, C'	33 conditions sur le partage d'un graphe simple orienté en deux
$\overset{sc}{\curvearrowright}$	37 chemin bigarré menant à une couleur scindante
$\overset{b}{\curvearrowright}$	37 chemin bigarré menant à un cycle bigarré
n -connectés	37 reliés par un chemin bigarré
$\alpha^+[g]$	37 une des deux parties d'un graphe complet biparti α
$\alpha[g]$	37 les sommets du sous-graphe de couleur α

$G - \alpha$	37	G privé de ses arêtes de couleur α
$\mathcal{F}$	56	famille d'arbres de sous-formules
$\bullet_i$	57,94	sommets d'un réseau correspondant à des coupures
Cut	57,94	ensembles des liens coupures d'un réseau
Ax	57,94	ensemble des liens axiomes d'un réseau
$\perp^\circ$	78	antiphases
1°	78	orthogonal des antiphases
$\#(\pi)$	79	indice d'une preuve séquentielle
$\#(\Pi)$	80	indice d'un réseau
$ff(\Pi)$	80	nombre de lien $\perp$ d'un réseau
$cc(\Pi)$	80	nombre de composantes connexes de tout sous graphe bigarré maximal d'un réseau
Π_∇	66	traduction du réseau Π en agrégat suivant ses liens par
$\Pi_\otimes$	68,127	traduction du réseau Π en agrégat suivant ses liens tenseur
$ A $	87	trame de l'espace cohérent A
$x \odot y[A]$	87	x et y sont deux points cohérents de la trame de A
$(\times \cap \cup)$	87	(incohérents, strictement cohérents, strictement incohérents)
$cl(X)$	87	clique d'un espace cohérents
$\prod_i A_i$	90	produit ordonné des espaces cohérents A_i ordonnés par i
$X \text{---} Y$	93	arête de sommets X et Y
$X \rightarrow Y$	93	arc joignant X et Y
$(\mathcal{F}, A)$	93	famille d'arbres de sous-formules orientés
F^∇	110	module canoniquement associé à une formule F
$i : \sim$	113	les points de $ A_i $ sont strictement cohérents, etc..
T^Π	129	condition correspondant à T dans un graphe

Index

Agregat, 35
orienté 127

Arbre des sous-formules, 56, 93
bien coloré, 57, 94

Bien coloré, arbres des sous-formules, 57, 94

Bigarré, 37, 59, 97
sous graphe bigarré maximal, 50

Contractile, ordre, 22

Contraction d'un ordre suivant deux points
équivalents, 18
équivalents-inférieurs, 19

Critère de correction des réseaux
avec mélange, 59
avec neutres et mélange, 80
ordonnés, 97

Cohérent, espace, 87

B-Connectés, points d'un graphe coloré, 37

Connecteurs généralisés, 91, 93, 111
définissables, 106

Elimination des coupures,
avec mélange, 56
avec neutres et mélange, 80
ordonnées, 98, 122

Entrelacs d'un réseau orienté, 127

Équivalents, points d'un ordre, 17

Équivalents-inférieurs, points d'un ordre, 18

Étirement d'un agrégat, 37

η -Expansion,
avec neutres et mélange, 83
ordonnée 104

Expansion d'un ordre suivant deux points
équivalents, 21
équivalents-inférieurs, 21

Expérience d'un réseau, 100

Extension orthogonale de deux ordres,
disjoints, 29
ayant un point commun, 30

Feuille
d'un (pré)réseau, 57, 94
d'un agrégat, 49
d'un agrégat orienté, 127

Frontière, 128
séparable, 128
séparée, 128

Frontière tendue, 129
séparable, 129
séparée, 129

Graphe
simple, 32
simple orienté, 93

Graphe de correction d'un réseau orienté, 97

Héréditairement scindant, tenseur, 75

Indice,
d'une preuve séquentielle avec neutres et
mélange, 79
d'un réseau avec neutres et mélange, 80

Lien d'un réseau
avec mélange, 58
avec neutres et mélange, 80
ordonné, 95

Modules d'un réseau ordonné, 103

Orthogonalité des relations,
faible, 27
forte, 28

Phases, sémantique des, 78

Parcours d'un réseau ordonné, 97
interne, 99

Précède,
espace cohérent, 89
lien, 96
règle séquentielle, 114, 122

Préseau,
avec mélange, 57
avec mélange et neutres, 80
ordonné, 94

Produit ordonné d'espaces cohérents, 91

Réseau,
avec mélange, 59
avec mélange et neutres, 80
ordonné, 97

Satellite d'un réseau ordonné, 127

Saturation d'un réseau ordonné, 134

Scindant(e),
couleur, 37
par , 64
tenseur , 64, 126

Sémantique,
des phases, 78
cohérente, 87, 100

Séquentialisation,
avec mélange, 63
avec mélange et neutres, 81
ordonnée, 125

Tendu(e),
réseau, 126
frontière, 129

Trame d'un espace cohérent, 87

Table des matières

Introduction	1
§A. Résumé	5
§B. Interprétation Naïve du Calcul Ordonné comme Modèle du Parallélisme	12
§C. Terminologie et Notations	14
 A Combinatoire	 15
1 Ordres et Orthogonalité	17
§A. Points Equivalents et Equivalents-Inférieurs	17
§B. Contractions d'un Ordre	18
§C. Expansions d'un Ordre	20
§D. Ordres Contractiles	22
§E. Orthogonalité des Relations Binaires	27
§F. Ordres et Orthogonalité	29
§G. Circuits dans un Graphe et Orthogonalité	31
 2 Agrégats	 35
§A. Introduction	35
§B. Agrégats	36
§C. Enoncé des Résultats	37
§D. Etirement d'un Agrégat	38
¶D.1. Etirement et Couleurs Scindantes	40
¶D.2. Etirement et Chemins Bigarrés	42
¶D.3. Etirement et Configurations Recherchées.	44
§E. Démonstration du Résultat Principal	46
¶E.1. L'Algorithme	46
¶E.2. Correction	47
¶E.3. Complexité de l'Algorithme en l'Absence de Cycle Bigarré	48
§F. Corollaires	49
§G. Sous-Graphes Bigarrés Maximaux	50

B	La règle de mélange	53
3	Séquents, Réseaux et Règle de Mélange	55
§A.	Le Calcul des Séquents	55
§B.	Les Réseaux Correspondants	56
§C.	Des Preuves Séquentielles aux Réseaux	59
4	Des Réseaux aux Séquents	63
§A.	Le théorème de séquentialisation	63
§B.	Réseaux et Agrégats	66
¶B.1.	Existence d'un par Scindant	66
¶B.2.	Existence d'un tenseur Scindant	68
¶B.3.	Optimisation de la méthode du tenseur scindant	73
¶B.4.	Agrégats et réseaux	73
¶B.5.	Existence d'un tenseur héréditairement scindant	75
5	Eléments neutres	77
§A.	Le Calcul des Séquents avec Mélange et Eléments Neutres	78
¶A.1.	Une Sémantique des Phases	78
¶A.2.	Les Règles du Calcul des Séquents	78
§B.	Les Réseaux Correspondants	80
¶B.1.	Elimination des Coupures	80
§C.	Séquentialisation	81
§D.	η -Expansion	83
C	Réseaux et séquents ordonnés	85
6	Des espaces cohérents au connecteur "précède"	87
§A.	Espaces cohérents	87
§B.	Le connecteur précède	88
§C.	Produit ordonné d'espaces cohérents	90
§D.	Sémantique cohérente du calcul multiplicatif avec mélange	92
7	Réseaux ordonnés	93
§A.	Définition des réseaux ordonnés	93
§B.	Elimination des Coupures Ordonnées	97
¶B.1.	Coupure sur un axiome [az]	98
¶B.2.	Coupure entre deux précède [pcd]	98

¶B.3. Coupure entre un tenseur et un par [ts/par]	99
¶B.4. Normalisation forte et confluence du calcul des réseaux ordonnés	100
§C. Sémantique cohérente des réseaux ordonnés	100
§D. Modularité de ce Calcul	103
§E. η -expansion	104
§F. Rapport avec le calcul de la deuxième partie.	104
§G. Connecteurs n-aires définissables	106
8 Le Calcul des Séquents Ordonnés	111
§A. Présentation	111
§B. Règles à une prémisses	112
¶B.1. L'axiome	112
¶B.2. La règle d'entropie	113
¶B.3. Le par ordonné	113
3.a. Sémantique cohérente de la règle par	113
¶B.4. Le connecteur précède	114
4.a. Sémantique cohérente de la règle du précède	114
§C. Règles à deux prémisses	115
¶C.1. La règle de MÉLANGE	115
1.a. Sémantique cohérente de la règle de MÉLANGE	116
¶C.2. Le TENSEUR ordonné	117
2.a. Sémantique cohérente de la règle du TENSEUR	118
¶C.3. La règle de COUPURE	120
¶C.4. Sémantique cohérente de la règle COUPURE	120
§D. Quelques Propriétés de ce calcul	122
¶D.1. Toute règle précède à deux prémisses est simulable	122
¶D.2. Non commutation du TENSEUR et du MÉLANGE	122
¶D.3. Elimination des coupures	122
9 Séquentialisation	125
§A. Liens par et précède finaux	125
§B. Structure des réseaux sans par ni précède final	126
§C. Monstres Choisis	130
§D. Séquentialisation	133
Conclusion	139
Bibliographie	140

Index des notations	143
Index	145

