
UNIVERSITÉ PARIS VII — DENIS DIDEROT

THÈSE DE DOCTORAT — MATHÉMATIQUES

Réseaux, cohérence et expériences obsessionnelles

Lorenzo Tortora de Falco

Soutenue le 28 janvier 2000

Directeur de thèse :
Vincent Danos

Rapporteurs :
Michele Abrusci
Simone Martini

Jury :
Michele Abrusci
Pierre-Louis Curien
Vincent Danos
Thomas Ehrhard
Jean-Yves Girard (Président)
Mario Girardi
Jean-Louis Krivine
Laurent Regnier
Jacqueline Vauzeilles

Alle sdusde, ed al loro coraggio.

«“Forse tu non pensavi ch'io loico fossi” !»

Dante, Inf., XXVII, 122-123

Roma, 10 ottobre 1991

Paris, 28 janvier 2000

Remerciements

Quiconque ait eu l'occasion de séjourner un tant soit peu au sein de l'équipe de logique mathématique de Paris 7, connaît l'extraordinaire potentiel que recèle un tel environnement scientifique. J'ai eu la chance d'y évoluer pendant plusieurs années, et cette thèse n'est que l'un des produits que j'ai pu en tirer. C'est avant tout à l'ensemble de cette équipe et à tous ceux qui ont contribué à son existence, que j'adresse mes remerciements.

Vincent Danos a dirigé mes recherches en me laissant une très grande liberté, tout en m'accordant l'attention dont j'avais besoin. Je crois que c'est une attitude non commune chez un directeur de thèse, de laquelle j'ai certainement beaucoup appris. J'ai aussi bénéficié de l'*impressionnante* vivacité de son esprit, de sa curiosité intellectuelle, de son enthousiasme, de son énergie. C'est grâce à sa compétence que j'ai pris connaissance des problèmes ouverts et ai commencé à porter mon propre regard sur les questions abordées dans la thèse.

Je remercie grandement Jean-Baptiste Joinet et Harold Schellinx, qui ont suivi mes travaux depuis le début.

J'ai toujours trouvé chez Jean-Baptiste une écoute attentive et des suggestions pertinentes.

Harold est le seul à avoir eu le courage de lire le manuscrit dans sa toute première version : si celui-ci est à présent un peu plus lisible, c'est en bonne partie grâce à ses commentaires.

Pas une seule ligne de ce qui suit ne serait écrite sans les contributions de Jean-Yves Girard à la théorie de la démonstration. Ses travaux ont constitué et constituent toujours pour moi une source inépuisable d'idées, de techniques, de réflexions, d'une telle richesse et variété que je ne cesse de m'en étonner. Je le remercie d'avoir bien voulu être président de mon jury, et du soutien qu'il ne m'a jamais ménagé.

Je remercie Jean-Louis Krivine ...d'exister ! L'élégance, l'originalité, la rigueur et l'exigence de sa démarche scientifique ont été pour moi l'un des

grands enseignements de ces années de recherche. Je le remercie aussi pour son soutien, et pour avoir accepté que ce travail débute sous sa responsabilité. C'est un très grand honneur qu'il me fait en acceptant de faire partie du jury.

Pour son soutien et pour son amitié, je remercie Marco Pedicini à qui la vie m'a très profondément lié à travers des chemins mystérieux . . .

Je remercie Myriam Quatrini pour son amitié, et pour tous les (nombreux) accueils, toujours chaleureux, qu'elle m'a réservés à Marseille. Le chapitre 15 porte la marque de sa collaboration et de celle (plus tardive) d'Olivier Laurent.

Merci aussi à tous les membres de l'équipe de Logique de la Programmation, pour la sympathie qu'ils m'ont témoignée lors de mes passages à Marseille.

Cette thèse a été en partie subventionnée par une bourse du Consiglio Nazionale delle Ricerche (C.N.R.). Je remercie à ce sujet Carlo Scoppola, pour ses conseils, ses nombreuses suggestions et son soutien.

Le soutien de François Métayer m'a valu un poste d'ATER à Nanterre. Cela m'a permis de terminer ma thèse, et surtout, de bénéficier d'une expérience d'enseignement très enrichissante. Je le remercie aussi pour les nombreuses (et toujours très agréables) discussions, non seulement mathématiques, que nous avons eues.

Je remercie Michele Abrusci de s'être intéressé à moi et à mon travail, et d'avoir aussi accepté de servir de rapporteur pour cette thèse. Tous mes vœux pour que ce soit le début d'une fructueuse interaction.

Je remercie Simone Martini d'avoir accepté de rapporter sur mon travail, et Pierre-Louis Curien, Mario Girardi, Thomas Ehrhard, Laurent Regnier et Jacqueline Vauzeilles d'avoir accepté de faire partie du jury.

Un merci chaleureux à Paul Ruet, pour avoir bien voulu partager avec moi un août romain particulièrement difficile.

Je remercie tous les thésards du département de mathématiques de Paris 7, pour les discussions de toutes sortes que nous avons eues, et pour leurs conseils de LATEX . . .

Impossible d'oublier la générosité d'Astyanax (alias Vincent Astier) et de Benoît Mariou, qui ont bien voulu faire cours à ma place dans la période de "débordement maximum".

Enfin, une pensée toute spéciale pour celle qui fut longtemps ma voisine de bureau, pour son sourire et sa spontanéité.

La qualité du travail administratif de Michèle Wasse est reconnue par tous les membres du département de maths. Mais je tiens à la remercier plus

particulièrement pour l'attention (extraordinaire) qu'elle accorde à chacun de ses étudiants.

Merci aussi à Claude Orioux pour sa gentillesse et sa disponibilité.

Un clin d'oeil amical à Odile Ainardi, qui sait faire régner une ambiance si agréable dans le secrétariat de l'équipe de logique. Et au prochain tour de valse !

Je suis très reconnaissant envers René Cori pour son accueil et sa disponibilité lors de mon arrivée à Paris. Il m'a, par ailleurs, sauvé la vie la nuit du 29 octobre 1999.

Je dois à Gianfranco Mascari de m'avoir initié à la théorie de la démonstration et de m'avoir introduit dans le milieu scientifique international. Et ça ne s'oublie pas.

Introduction

Le sujet de notre dissertation le jour du baccalauréat de philosophie, était :
“Qu’est-ce qu’une preuve ?”

Nous n’osons imaginer ce que serait la réaction de notre correcteur d’alors devant le texte qui suit... Pourtant, celui-ci peut bien être lu comme une tentative de donner des éléments de réponse (mathématiques) à cette question.

Pour ce faire, il va falloir lui donner un sens (un tant soit peu) plus précis. S’il prend le parti (comme nous le ferons) de s’intéresser seulement aux preuves mathématiques, le lecteur “méditant” la suite de signes écrits sur une feuille de papier que les mathématiciens s’accordent tous pour appeler une preuve, se rendra compte rapidement que plusieurs “représentations” de cette preuve sont possibles. Mais il n’aura pas vraiment avancé dans la réflexion sur la question posée initialement, faute de *moyens* pour l’attaquer.

L’élimination des coupures

Au début des années 30, Gentzen démontre un théorème célèbre, connu sous le nom de “théorème d’élimination des coupures”. Une façon (*très*) informelle d’en exprimer le contenu pourrait être : “toute preuve courte, intelligente et compréhensible d’un énoncé A , peut être transformée en une preuve longue, idiote et incompréhensible de A ”.

Palpitant. Oui, palpitant en vérité, puisque la transformation de la preuve intelligente de A en la preuve idiote de A (le processus d’élimination des coupures) correspond (en un sens très précis) à un processus de *calcul*.

Voici donc un premier moyen pouvant nourrir notre réflexion : si une preuve cache un processus de calcul, le comprendre sera une façon de comprendre la preuve à laquelle il est sous-jacent.

La démonstration de Gentzen suggère de penser à la règle de coupure (la seule règle “intelligente”, dont le théorème en question prouve l’inutilité en

logique pure) comme à un moyen dont disposent les preuves pour communiquer. L'interaction qui en découle est le calcul qui est associé à l'élimination de la coupure. Ceci conduit à un changement de point de vue : au lieu de "contempler" les preuves, on va essayer d'en décrire les interactions possibles avec les autres preuves, et ce en termes de calcul.

L'objet de notre étude sera la structure de cette interaction, l'objectif à terme étant de rendre compte des démonstrations mathématiques comme d'objets géométriques qui en expriment le contenu calculatoire.

L'intérêt d'une telle étude n'est pas seulement logique. Dans les années 60, la correspondance entre preuves et programmes est clairement établie par l'isomorphisme de Curry-Howard : une preuve de logique intuitionniste est un programme, et l'exécution de ce programme correspond très exactement à l'élimination des coupures dans la preuve de logique intuitionniste.

Du point de vue informatique, cette correspondance suggère une nouvelle méthode de programmation, dite de "programmation par preuves" (cf. [KriPar 90]).

Ayant la possibilité de désigner et définir précisément les programmes et les données au niveau logique, on peut programmer en restant au niveau logique (en ne manipulant que des démonstrations). Très schématiquement, pour programmer, par exemple, la fonction récursive totale f , on procède de la façon suivante :

- on donne un système d'équations (Eq) définissant f
- on exhibe une preuve formelle π (dans un système de déduction donné) de $\forall x(Nx \rightarrow Nfx)$ (c.à.d. une preuve de "pour tout x , si x est un entier alors $f(x)$ l'est aussi")
- on associe (par l'isomorphisme de Curry-Howard) à π le programme t_f .

La puissance de cette approche de la programmation réside précisément dans la possibilité de programmer en ne manipulant que des objets mathématiques. Ceci permet notamment de prouver *abstraitement* que :

- quelque soit la stratégie utilisée pour exécuter un programme associé à une preuve, l'exécution termine (théorème de normalisation forte)
- deux stratégies différentes d'exécution d'un programme associé à une preuve produisent le même résultat (théorème de confluence)
- tout programme associé à une preuve est correct, c.à.d. que, dans l'exemple ci-dessus, le programme t_f calcule bien les valeurs de la fonction f spécifiée par (Eq).

Les deux approches traditionnelles dans l'étude des phénomènes calculatoires (non seulement ceux sous-jacents aux preuves qui nous occuperont dans cette thèse) sont connues sous les noms de "syntaxe" et de "sémantique".

La première est une description que l'on qualifie d'“intentionnelle” des objets étudiés, tandis que la deuxième en est une description “extensionnelle”. La syntaxe est très précise dans sa description des preuves, mais elle est souvent redondante et manque de clarté mathématique (les objets syntaxiques n'étant pas, en général, des objets mathématiques usuels). La sémantique, au contraire, rend compte des démonstrations comme de fonctions, mais nous fait payer le prix de cette simplicité par son imprécision et son incomplétude dans la représentation des preuves qu'elle fournit.

Nous entamerons notre étude en faisant usage de ces deux approches.

La logique linéaire

C'est la notion de stabilité, découverte par Berry en 1978 ([Berry 78]), qui conduit Jean-Yves Girard à proposer, dans les années 80, les espaces cohérents comme structures pour interpréter les formules de la logique intuitionniste. Girard s'aperçoit alors que les connecteurs intuitionnistes peuvent être (mathématiquement) décomposés. Il n'était pas évident à priori que cette décomposition pouvait être internalisée, c.à.d. exprimée au moyen de nouvelles opérations logiques. Girard montre que ceci est en fait possible : par exemple l'implication intuitionniste $A \rightarrow B$ peut s'écrire comme $!A \multimap B$, où “ $\multimap$ ” désigne l'implication linéaire qui a le sens d'une causalité (algorithmiquement, $\sigma \multimap \tau$ est le type des algorithmes fonctionnels de σ vers τ qui appellent leur argument exactement une fois) et “ $!$ ” a le sens d'une répétition “ A autant de fois que l'on veut” et correspond algorithmiquement à une mise en mémoire. C'est ainsi que, dans [Girard 87], apparaît la logique linéaire, un raffinement de la logique intuitionniste et de la logique classique qui se caractérise par la décomposition des connecteurs logiques usuels en deux classes (les multiplicatifs et les additifs) et par l'introduction de nouveaux connecteurs (les exponentielles) qui permettent une gestion logique des règles structurelles d'affaiblissement et de contraction.

Ce changement de point de vue a eu des conséquences très frappantes en théorie de la démonstration, une des plus remarquables étant l'introduction des réseaux de preuves, qui fournissent une représentation géométrique du calcul.

Contenu de la thèse

La représentation géométrique des démonstrations fournie par les réseaux de preuves est certes imparfaite, partielle, souvent lourde et insatisfaisante,

mais elle a cependant l'immense vertu d'exister. Nous en ferons le point de départ de notre analyse.

Nous passerons ensuite à une étude plus sémantique des preuves (toujours représentées sous formes de réseaux), en utilisant des structures traditionnelles et rudimentaires telles que les ensembles et les espaces cohérents. Il ne s'agira pas de faire des exploits dans la manipulation d'objets mathématiques raffinés avec plein de propriétés, mais plutôt de chercher à établir des liens simples et clairs entre différentes représentations des preuves, avec l'espoir d'apprendre quelque chose de nouveau sur celles-ci.

Nous discuterons enfin de ce que la logique linéaire a pu nous apprendre sur les preuves de la logique classique (celle-ci étant la logique de la pratique mathématique courante).

Plus précisément, cette thèse s'articule en trois parties :

- I) Une étude syntaxique des réseaux de preuves : chapitres 1-5 (et annexe A).
- II) Une étude sémantique des réseaux de preuves : chapitres 6-14.
- III) Une étude syntaxique et sémantique des preuves de la logique classique : chapitre 15 (et annexe B).

Le contenu des 4 premiers chapitres est essentiellement celui du papier [Torto 00a]. Le résultat principal est le théorème de confluence faible (le théorème 4.3.5). Nous pensions pouvoir le démontrer en une trentaine de pages, mais ce ne fut pas (hélas) le cas. Pour le présenter dans toute sa généralité, il fallait avoir à la disposition la syntaxe des réseaux pour toute la logique linéaire (LL). Quelle fut notre surprise en apprenant que la *seule* référence en la matière était le premier papier de Girard, [Girard 87]. Mais il était impossible de ne pas tenir compte des innovations intervenues depuis, et nous avons donc pris la peine de faire le point sur les réseaux de la logique linéaire.

Nous avons *tout* démontré (ou redémontré) sur la correction des réseaux et sur leur normalisation dans les chapitres 1 et 2. En l'absence de connecteurs additifs, il fallait adapter la gestion des quantificateurs de [Girard 91a] aux réseaux de MELL (le fragment multiplicatif et exponentiel de LL), définis dans [Danos 90]. Nous avons beaucoup bénéficié, dans cette oeuvre, des notes prises lors d'un cours de DEA tenu par Vincent Danos. Il s'agissait ensuite de gérer la présence des connecteurs additifs de façon compatible avec cette notion de réseau. Nous insistons sur le fait que si plusieurs idées (notamment l'introduction des "sauts") étaient connues, elles n'apparaissent qu'en appendice de quelques papiers : nulle part il n'en est fait (à notre connaissance) un usage précis. L'ambition de ces deux premiers chapitres est de

constituer un texte qui puisse servir de référence sur le sujet.

Dans le chapitre 3 nous prouvons une propriété spécifique à la normalisation en présence d’additifs (le théorème de séparation). C’est une propriété de la duplication qui accompagne l’étape élémentaire commutative additive, et qui est l’un des principaux ingrédients de la preuve du théorème de standardisation additive de l’annexe A (le théorème A.2.2).

Le chapitre 4 est consacré au théorème de confluence faible : nous montrons que, même en présence des additifs, la normalisation des réseaux (dans toute sa généralité) est confluente, pourvu que les conclusions des réseaux que l’on considère ne contiennent pas d’additifs (le théorème 4.3.5).

Dans l’annexe A, nous nous intéressons à la propriété de forte normalisation des réseaux. Les connaissances acquises dans les 4 premiers chapitres nous permettent de prouver le théorème de standardisation additive. Malheureusement ce résultat est encore insuffisant pour donner une preuve complète de forte normalisation des réseaux, problème qui est encore ouvert à ce jour (voir à ce sujet l’introduction de l’annexe A, et la remarque A.5.2). Le contenu de l’annexe A est celui du papier en cours de soumission [Torto 00b].

Le dernier chapitre de la première partie est consacré à une modification de la syntaxe des réseaux suggérée par la sémantique : nous introduisons les “multiboîtes additives” et les modifications dans les étapes de normalisation que ces nouveaux objets induisent. Ce chapitre est une contribution à l’étude de la normalisation en présence d’additifs, qui est notoirement un problème plutôt épineux. En nous plaçant dans le cadre de la logique linéaire multiplicative et additive, nous obtenons un théorème de forte normalisation (le théorème 5.3.8), nous montrons que la propriété de séparation du chapitre 3 est toujours satisfaite en présence de multiboîtes (proposition 5.3.4), et nous définissons une procédure de normalisation confluente (théorème 5.5.15).

Dans la deuxième partie, nous posons une question très naturelle que nous appelons “injectivité de la sémantique”. Il s’agit de comparer deux relations d’équivalence définies sur les réseaux : l’équivalence induite par le processus de calcul (syntaxique) et l’équivalence définie par la représentation “algébrique” de ce même calcul (sémantique). Nous utilisons, pour effectuer cette analyse, la notion d’expérience, définie dans [Girard 87] et rarement reprise depuis. On fera référence à trois types de sémantique qui s’adaptent bien à cette notion : la sémantique relationnelle (multiensembliste), la sémantique cohérente ensembliste et la sémantique cohérente multiensembliste.

Nous donnons les définitions de base dans le chapitre 6, nous posons clairement le problème dans le chapitre 7 et nous le résolvons dans un cas très

simple. En suivant la méthode de résolution du chapitre 7, nous définissons, dans le chapitre 8, la notion-clef d’“expérience obsessionnelle”, que nous étudions de façon approfondie (chapitres 8,9). A l’égard de cette notion, la sémantique cohérente a un comportement particulier : son uniformité permet, dans une certaine mesure, de déduire le caractère obsessionnel d’une expérience à partir de son résultat (propositions 9.3.4 et 9.3.5).

Nous exposons des résultats partiels dans le chapitre 10.

Les chapitres 11 et 12 permettent de ramener la question de l’injectivité de la sémantique cohérente à celle de l’existence d’un certain type d’expérience (une expérience “injective”) pour des réseaux très simples.

Au chapitre 13, coup de théâtre : une telle expérience, en général, n’existe pas. On en déduit immédiatement un contre-exemple à l’injectivité de la sémantique cohérente (ensembliste et multiensembliste).

Dans le chapitre 14 nous allons au bout de notre méthode : nous prouvons l’injectivité pour des fragments significatifs de LL (théorème 14.2.7, corollaire 14.2.8) et nous mettons en évidence l’existence d’un lien entre la connexité des réseaux et la propriété “d’uniformité” de la sémantique cohérente. Nous conjecturons que celui-ci est *très* fort.

La proposition 6.3.9 permet d’affirmer que tous les résultats d’injectivité établis dans ce dernier chapitre pour la sémantique cohérente multiensembliste, restent vrais pour la sémantique relationnelle.

Dans la troisième partie, nous étudions les liens entre la logique linéaire polarisée et la logique classique. Le contenu de ce chapitre est le fruit d’une collaboration en cours avec Myriam Quatrini et Olivier Laurent. Nous définissons une contrainte sur les preuves de la logique classique (la ρ -contrainte) qui, avec la η -contrainte de [DJS 97], fournit un système $(LK_{pol}^{\eta,\rho})$ complet pour la prouvabilité classique (théorème 15.4.9) et stable par élimination des coupures (théorème 15.4.11), dont l’image dans LL est la logique linéaire polarisée (proposition 15.4.12). C’est dans ce système que la question de l’injectivité de la sémantique cohérente pour la logique classique semble pouvoir trouver une réponse positive (remarque 15.5.4). La suite naturelle de ce chapitre est l’annexe B, dont le contenu recoupe celui d’un papier déjà publié ([Torto 97]).

Nous concluons, au chapitre 16, par une courte discussion au sujet des résultats (positifs et négatifs) de la thèse.

Terminons cette introduction en signalant que dans un domaine relativement nouveau comme celui qui nous occupe dans cette thèse, les problèmes rédactionnels sont importants. Nous avons tenté d’être compréhensibles tout en gardant un degré raisonnable de rigueur et de précision. Nous avons été

contraints pour cela à quelques lourdeurs, pour lesquelles nous nous excusons auprès du lecteur.

Table des matières

1 Réseaux avec sauts de PN_2	21
1.1 Formules de la logique linéaire et calcul des séquents	21
1.2 Pré-structures de preuve, pré-réseaux, réseaux	24
1.3 Les réseaux avec sauts	29
2 Normalisation des réseaux de PN_2	39
2.1 La notion de substitution	39
2.2 La normalisation	44
2.3 Préservation de la correction	52
3 La propriété de séparation	63
3.1 Le lemme de substitution	63
3.2 Normalisation et séparation	64
4 Confluence faible des réseaux de PN_2	69
4.1 La réduction $\neg[ccad]$	71
4.2 La réduction paresseuse	75
4.3 Le théorème de confluence faible	77
5 Les multiboîtes additives	83
5.1 De la non confluence aux multiboîtes	84
5.2 La normalisation des réseaux de $MALL$	87
5.3 Normalisation forte	92
5.4 Structure des réseaux de $MALL$	97
5.5 Procédure de normalisation confluente	101
5.6 Remarque sur le calcul HC de Girard	109
6 Sémantique des réseaux	111
6.1 Préliminaires, conventions, notations	111
6.2 Les expériences	114

6.3	Quelques propriétés des expériences	119
7	La question de l'injectivité	125
7.1	Position du problème	125
7.2	Sous-graphes d'un réseau	128
7.3	L'injectivité pour MLL	130
8	Les expériences obsessionnelles	133
8.1	Définition et propriétés	133
8.2	Reconstruction des graphes partiels	138
8.3	Reconstruction des graphes	141
9	Résultats obsessionnels	151
9.1	Le cas général	151
9.2	Le cas de la sémantique relationnelle	156
9.3	Le cas de la sémantique cohérente	158
10	Bilan sur l'injectivité : résultats et conjectures	163
10.1	Bilan pour la sémantique relationnelle	164
10.2	Bilan pour la sémantique cohérente	166
11	Elimination des boîtes	171
11.1	Réduction au cas des 1-expériences	171
11.2	Réduction au cas sans boîtes	176
12	Elimination des liens par	179
13	La sémantique cohérente n'est pas injective	185
13.1	Le premier contre-exemple	185
13.2	Le deuxième contre-exemple	188
13.3	Remarques au sujet de la non injectivité	191
14	Fragments d'injectivité	195
14.1	Le cas des contractions terminales	195
14.2	Conclusions	203
15	Logique linéaire polarisée et logique classique	209
15.1	Le calcul des séquents classique LK	210
15.2	La logique linéaire polarisée et focalisée LL_p^f	212
15.3	LK et LL	216
15.4	Le système $LK_{pol}^{\eta,\rho}$	221

15.5	Isomorphismes sémantiques et syntaxiques	230
16	Discussion : syntaxe et sémantique	241
A	Le théorème de standardisation additive	249
A.1	Normalization without interaction	251
A.2	The standardization theorem	265
A.3	The first projection of the reducts of T	269
A.4	The second projection of the reducts of T	272
A.5	Proof of the standardization theorem	283
B	Denotational semantics for polarized (but non-constrained)	
	LK by means of the additives	287
B.1	Introduction	287
B.2	Preliminaries	289
B.3	LK_{pol} and the embedding P^a in linear logic	291
B.4	Canonical form	296
B.5	Cut-elimination in LK_{pol}	300
B.6	Computational isomorphisms in LK_{pol}	303
B.7	Further développements	304

Chapitre 1

Réseaux avec sauts de PN_2

Ce chapitre est consacré à la définition des réseaux de preuves pour la logique linéaire du second ordre, introduits dans [Girard 87]. La variante que nous présentons ici tient compte des innovations intervenues depuis, et notamment de [Girard 91a] et [Danos 90]. La gestion des affaiblissements se fait au moyen d’une “fonction de saut”, dont la définition doit tenir compte de la présence des connecteurs additifs. Le chapitre se termine par l’incontournable théorème de séquentialisation (le théorème 1.3.10).

1.1 Formules de la logique linéaire et calcul des séquents

Nous allons définir les formules de la logique linéaire du second ordre, *sans* quantificateurs du premier ordre. Il n’y a absolument aucune difficulté dans l’extension de tout ce que nous allons dire au premier ordre. Par contre, omettre la quantification du premier ordre va nous permettre d’éliminer des redondances inutiles dans la rédaction.

1.1.1. DÉFINITION. Les connecteurs linéaires se répartissent en quatre groupes :

- le connecteur de la négation linéaire $\perp$
- les connecteurs multiplicatifs binaires $\otimes$ et $\wp$, qu’on appelle respectivement **tenseur** et **par**
- les connecteurs additifs binaires $\&$ et $\oplus$ qu’on appelle respectivement **avec** et **plus**
- les connecteurs exponentiels unaires $!$ et $?$ qu’on appelle respectivement **bien sûr** et **pourquoi pas**.

A ces connecteurs on rajoute les quantificateurs du second ordre $\forall$ et $\exists$.

On se donne un ensemble dénombrable de variables propositionnelles, et c'est à partir de ce langage que nous allons construire les formules de la logique linéaire (LL).

1.1.2. DÉFINITION. Les formules de la logique linéaire sont le quotient des formules bâties sur les connecteurs et les quantificateurs introduits en 1.1.1 par les équations suivantes (appelées lois de de Morgan) :

$$A^{\perp\perp} = A$$

$$\begin{aligned} (A \otimes B)^\perp &= A^\perp \wp B^\perp & (A \wp B)^\perp &= A^\perp \otimes B^\perp \\ (!A)^\perp &=?A^\perp & (?A)^\perp &=!A^\perp \\ (\exists X A)^\perp &= \forall X A^\perp & (\forall X A)^\perp &= \exists X A^\perp \end{aligned}$$

et par les équations égalisant les formules qui ne diffèrent que par le nom de leurs variables liées.

Deux formules A et B telles que $A = B^\perp$ sont dites duales.

Par abus de langage, nous utiliserons souvent le terme “formule” là où l'expression “occurrence de formule” serait plus appropriée, et ce afin d'éviter des lourdeurs inutiles.

1.1.3. DÉFINITION. Soit E un ensemble et E^* le monoïde libre engendré par E . On définit la relation d'équivalence $\simeq$ sur E^* : pour tout $m, m' \in E$, $m \simeq m'$ ssi $\forall e \in E$, m et m' contiennent le même nombre d'occurrences de e .

On appelle **multiensemble** d'éléments de E tout élément de l'ensemble quotient $E^* / \simeq$. L'ensemble $E^* / \simeq$ est appelé le monoïde libre commutatif engendré par E .

Un **séquent** (de LL) est un multiensemble de formules (de LL), souvent précédé du symbole $\vdash$.

1.1.4 Les calculs des séquents avec hypothèses LL_2 et LL'_2

Dans ce paragraphe, Γ et Δ sont des multiensembles de formules de LL. La dernière règle ci-dessous est appelée *mix* : c'est la seule règle de LL'_2 qui n'est pas une règle de LL_2 . Nous ferons toujours référence au calcul LL_2 , sauf dans les premiers chapitres, au cours desquels nous ferons quelques (rares) allusions à LL'_2 .

1.1. FORMULES DE LA LOGIQUE LINÉAIRE ET CALCUL DES SÉQUENTS 23

Axiome et coupure :

$$(Ax) \quad \vdash A, A^\perp \quad (cut) \quad \frac{\vdash \Gamma, A \quad \vdash \Delta, A^\perp}{\vdash \Gamma, \Delta}$$

Règles logiques multiplicatives :

$$(\otimes) \quad \frac{\vdash \Gamma, A \quad \vdash \Delta, B}{\vdash \Gamma, \Delta, A \otimes B} \quad (\wp) \quad \frac{\vdash \Gamma, A, B}{\vdash \Gamma, A \wp B}$$

Règles logiques additives :

$$(\oplus) \quad \frac{\vdash \Gamma, A}{\vdash \Gamma, A \oplus B} \quad \frac{\vdash \Gamma, B}{\vdash \Gamma, A \oplus B} \\ (\&) \quad \frac{\vdash \Gamma, A \quad \vdash \Gamma, B}{\vdash \Gamma, A \& B}$$

Règles exponentielles structurelles :

$$(?W) \quad \frac{\vdash \Gamma}{\vdash \Gamma, ?A} \\ (?C) \quad \frac{\vdash \Gamma, ?A, ?A}{\vdash \Gamma, ?A}$$

Promotion ou règle bien sûr :

$$(!) \quad \frac{\vdash ?\Gamma, A}{\vdash ?\Gamma, !A}$$

Règle de déréluction ou pourquoi pas :

$$(?de) \quad \frac{\vdash \Gamma, A}{\vdash \Gamma, ?A}$$

Règles pour les quantificateurs du second ordre (Y n'est pas libre dans Γ) :

$$(\forall_2) \quad \frac{\vdash \Gamma, A[Y/X]}{\vdash \Gamma, \forall X A} \quad (\exists_2) \quad \frac{\vdash \Delta, A[T/X]}{\vdash \Delta, \exists X A}$$

La règle Hypothèse :

$$(Hyp) \quad \overline{\vdash \Gamma}$$

La règle mix :

$$(mix) \quad \frac{\vdash \Gamma \quad \vdash \Delta}{\vdash \Gamma, \Delta}$$

1.2 Pré-structures de preuve, pré-réseaux, réseaux

1.2.1. DÉFINITION. (variante de [DanReg 95]) Une **pré-structure de preuve** est un graphe orienté dont les noeuds sont appelés liens, et dont les arêtes sont étiquetées par des formules de LL. Chaque lien est défini avec une arité et une coarité, c.à.d. avec un certain nombre d'arêtes incidentes appelées les prémisses du lien, et avec un certain nombre d'arêtes émergentes appelées les conclusions du lien. Les prémisses d'un lien sont aussi appelées les formules actives du lien.

- Un lien **hypothèse** ou H a $n \geq 1$ conclusion(s) chacune étiquetée par une formule, et aucune prémisses
- un lien **axiome** ou ax a deux conclusions étiquetées par deux formules duales et aucune prémisses
- un lien **coupure** ou cut a deux prémisses étiquetées par deux formules duales et aucune conclusion
- un lien **par** ou $\wp$ (resp. **tenseur** ou $\otimes$) a deux prémisses et une conclusion. Si la prémisses gauche du lien est étiquetée par la formule A et la prémisses droite du lien est étiquetée par la formule B , alors la conclusion du lien sera étiquetée par la formule $A \wp B$ (resp. $A \otimes B$)
- un lien **bien sûr** ou $!$ a une prémisses, et une conclusion étiquetée par le **bien sûr** de l'étiquette de la prémisses
- un lien **deréliction** ou $?de$ a une prémisses, et une conclusion étiquetée par le **pourquoi pas** de l'étiquette de la prémisses
- un lien **affaiblissement** ou $?w$ n'a aucune prémisses, et une conclusion étiquetée par $?A$, où A est une formule
- un lien **contraction** ou $?co$ a deux prémisses et une conclusion, toutes étiquetées par la même formule $?A$
- un lien **porte auxiliaire** ou **pax** a une prémisses et une conclusion, toutes deux étiquetées par la même formule $?A$
- un lien **1-plus** ou $\oplus_1$ (resp. **2-plus** ou $\oplus_2$) a une prémisses et une conclusion, telles que si A est l'étiquette de la prémisses, alors $A \oplus B$ (resp. $B \oplus A$) est l'étiquette de la conclusion, où B est une formule. Nous parlerons aussi d'un lien **plus** ou $\oplus$ lorsque nous ne voudrons pas préciser s'il s'agit d'un lien $\oplus_1$ ou $\oplus_2$.
- un lien **avec** ou $\&$ a deux prémisses et une conclusion. Si la prémisses gauche du lien est étiquetée par la formule A et la prémisses droite du lien est étiquetée par la formule B , alors la conclusion du lien sera étiquetée par la formule $A \& B$
- un lien **coad** a deux prémisses et une conclusion, toutes étiquetées par la même formule

- un lien **pour tout** ou $\forall$ (resp. **il existe** ou $\exists$) a une prémisses et une conclusion, telles que si A (resp. $A[T/X]$) est l'étiquette de la prémisses, alors $\forall X A$ (resp. $\exists X A$) est l'étiquette de la conclusion. Nous dirons que la variable X est la variable propre du lien $\forall$.

Soit G un graphe orienté obtenu en utilisant les liens ci-dessus, et t.q. :

- (α) toute arête de G est conclusion d'un unique lien
- (β) toute arête de G est prémisses d'au plus un lien.

Les arêtes de G qui ne sont prémisses d'aucun lien sont les conclusions de G .

Nous dirons qu'un tel graphe G est une pré-structure de preuve s'il satisfait les trois conditions suivantes :

(1) boîte ! :

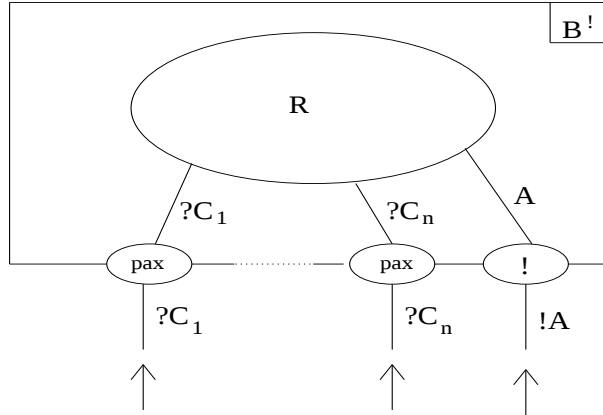
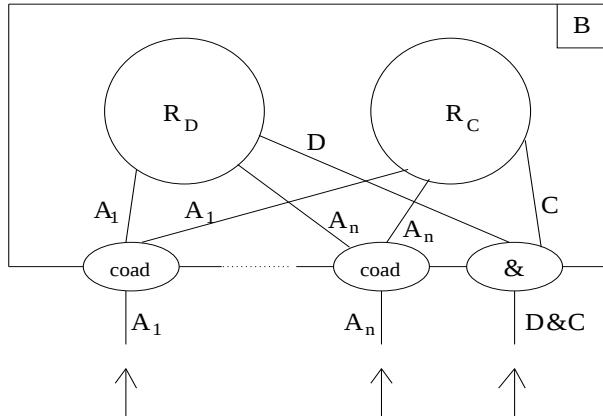
- (1.i) à tout lien bien sûr n est associé un (unique) sous-graphe orienté $B^!$ de G (satisfaisant (α) et (β)), tel que une parmi les conclusions de $B^!$ est la conclusion de n et toute autre conclusion de $B^!$ (il pourrait n'y en avoir aucune) est conclusion d'un lien pax. $B^!$ est appelé une boîte exponentielle et est représenté comme dans la figure 1 ; n est appelée la porte principale (pal en abrégé) de $B^!$
- (1.ii) à tout lien pax n est associée une boîte exponentielle $B^!$ de G t.q. une parmi les conclusions de $B^!$ soit la conclusion de n (voir fig.1) ; on dit que n est une porte auxiliaire de $B^!$

(2) boîte & :

- (2.i) à tout lien avec n est associé un (unique) sous-graphe B de G (satisfaisant (α) et (β)), tel que une parmi les conclusions de B est la conclusion de n et toute autre conclusion de B (il pourrait n'y en avoir aucune) est conclusion d'un lien coad. B est appelé une boîte additive et est représenté comme dans la figure 2 ; n est appelée la porte & de B
- (2.ii) à tout lien coad n est associée une boîte additive B , telle que une parmi les conclusions de B est conclusion de n (voir fig.2) ; on dit que n est une porte coad de B

(3) condition d'emboîtement :

deux boîtes quelconques (additives ou exponentielles), sont disjointes, ou bien l'une d'entre elles est contenue dans l'autre.

Figure 1. La boîte exponentielle $B^!$ de la pré-structure de preuve G .Figure 2. La boîte additive B de la pré-structure de preuve G .

Les liens suivants sont appelés **liens logiques** : par, tenseur, bien sûr, dérélction, 1-plus, 2-plus, avec, pour tout, il existe. Par, tenseur, avec, 1-plus et 2-plus sont aussi appelés **liens logiques propositionnels**. Les arêtes d'une pré-structure de preuve qui ne sont prémisses d'aucun lien, sont les conclusions de la pré-structure de preuve.

Les arêtes d'une pré-structure de preuve étant orientées "des axiomes vers les conclusions", il n'y aura pas d'ambiguïté lorsque nous dirons d'un chemin qu'il "monte" où qu'il "descend" dans la pré-structure de preuve.

Nous parlerons souvent de boîtes, liens ou arêtes d'une pré-structure de preuve R contenus dans une boîte B (additive ou exponentielle) de R . Dans le cas des liens, nous ne considérerons pas les portes de B comme des liens contenus dans B . En écrivant "un lien l (resp. une arête a) d'une boîte B (additive ou exponentielle)" d'une pré-structure de preuve donnée, nous sous-entendons toujours que l (resp. a) est contenu dans B ou bien que l (resp. a) est une porte (resp. une conclusion) de B . Si B est une boîte (additive ou exponentielle) d'une pré-structure de preuve R , alors la plus grande (resp. la plus petite) boîte de R contenant B est clairement définie, grâce à la condition d'emboîtement de la définition 1.2.1.

Nous dirons qu'un lien ou une arête d'une pré-structure de preuve R donnée est à **profondeur** exponentielle (resp. additive) n dans R , si il ou elle est contenu(e) dans exactement n boîtes exponentielles (resp. additives) de R . Dans le cas d'une boîte (additive ou exponentielle) B , nous dirons que B est à **profondeur** exponentielle (resp. additive) n dans R , si B est contenue dans exactement n boîtes exponentielles (resp. additives) de R , toutes différentes de B . Lorsque R est un réseau, cette même définition s'étend au cas d'un sous-réseau de R (défini dans 1.3.12).

1.2.2. REMARQUE. Soit B une boîte (additive ou exponentielle) de la pré-structure de preuve R . Une conséquence immédiate de la définition 1.2.1 est que le contenu de B est une pré-structure de preuve.

1.2.3. DÉFINITION. (pré-réseau) Soit R une pré-structure de preuve. On dira que R est une **structure de preuve** ou un **pré-réseau**, lorsque :

- (i) si une variable apparaît libre dans la formule étiquette d'une conclusion de R , alors ce n'est la variable propre d'aucun lien $\forall$ de R
- (ii) tout(e occurrence de) lien $\forall$ fait usage d'une variable propre différente
- (iii) si X est la variable propre d'un lien $\forall$ d'une boîte B de R , alors toute occurrence de X est contenue dans B
- (iv) pour toute boîte B de R , la pré-structure de preuve contenue dans B est un pré-réseau.

1.2.4. REMARQUE. Derrière les conditions de la définition 1.2.3, se cache en fait le (pénible) phénomène du renommage des variables : si par exemple on a deux pré-réseaux R_1 et R_2 et on souhaite connecter une conclusion de R_1 et une conclusion de R_2 par un lien $\otimes$, pour que les conditions de la définition soient satisfaites par le réseau ainsi obtenu on peut être amené

à renommer des variables apparaissant libres dans des formules de R_1 ou/et des variables apparaissant libres dans des formules de R_2 .

1.2.5. DÉFINITION. (graphe apparié 1) Nous dirons que deux arêtes d'un graphe orienté sont coïncidentes, lorsqu'elles ont le même but. On dit que le couple $(G, App(G))$ est un graphe **apparié** lorsque G est un graphe orienté et $App(G)$ est un ensemble de n -uples ($n \geq 2$) d'arêtes deux à deux distinctes et coïncidentes.

Soit R un pré-réseau et $B_1, \dots, B_k$ les boîtes de R de profondeur additive et exponentielle zéro. On va associer à R un ensemble $App(R)$ et un graphe apparié $R_{ap} = (G_R, App(R))$.

G_R est obtenu à partir de R de la façon suivante :

- on substitue à chaque boîte B_i avec p_i conclusions ($i \in \{1, \dots, k\}$) un noeud H avec p_i conclusions
- pour chaque lien $n \forall$ de variable propre X_n et pour chaque lien m qui n'est pas un lien de l'une des boîtes B_i et dont une prémisse a une étiquette contenant X_n libre, on rajoute une arête ayant comme source m et comme but n
- pour chaque lien $n \forall$ de variable propre X_n et pour chaque lien m qui est un lien de l'une des boîtes B_i et dont une prémisse a une étiquette contenant X_n libre, on rajoute une arête ayant comme source le noeud H qui substitue la boîte contenant m et ayant comme but n .

L'ensemble $App(R)$ contient les m -uples suivants (et eux seuls) :

- les couples de prémisses de chaque lien $\wp$ et $?co$ de profondeur additive et exponentielle zéro dans R
- les p -uples de toutes les arêtes incidentes dans un lien $\forall$ de G_R (parmi lesquelles il y a bien sûr toujours au moins la prémisse du lien $\forall$ considéré).

1.2.6. DÉFINITION. (graphe de correction 1) Soit R un pré-réseau et $B_1, \dots, B_k$ les boîtes de R de profondeur additive et exponentielle zéro. Soit $R_{ap} = (G_R, App(R))$ le graphe apparié associé à R par la définition 1.2.5.

Un **switching** S de R est le choix d'une arête pour chaque n -uple de $App(R)$. A tout switching S de R est associé un graphe non orienté $S(R)$, appelé **graphe de correction** : on efface, dans G_R , pour chaque n -uple de $App(R)$, toutes les arêtes qui ne sont pas sélectionnées par S . On oublie ensuite les étiquettes et l'orientation des arêtes. On notera $S(R)$ le graphe de correction de R associé à S .

1.2.7. DÉFINITION. (réseau -préliminaire-) Soit R un pré-réseau et $B_1, \dots, B_k$

les boîtes de R de profondeur additive et exponentielle zéro. On dira que R est un **réseau** lorsque les conditions suivantes sont satisfaites :

- R satisfait (AC) : pour tout switching S de R , le graphe de correction $S(R)$ est acyclique
- pour toute boîte exponentielle $B_i \in \{B_1, \dots, B_k\}$, le pré-réseau R_i contenu dans B_i est un réseau
- pour toute boîte additive $B_i \in \{B_1, \dots, B_k\}$, il existe deux réseaux disjoints R_i^1 et R_i^2 contenus dans B_i , t.q. tout lien contenu dans B_i est soit un lien de R_i^1 soit un lien de R_i^2 , toute conclusion de R_i^1 (resp. R_i^2) est prémisses d'une porte de B_i , et toute porte de B_i a deux prémisses : une conclusion de R_i^1 et une conclusion de R_i^2 .

Si B est une boîte additive du réseau T , alors (par définition de réseau) il existe un sous-pré-réseau R de T qui est un réseau et t.q. B est à profondeur exponentielle et additive zéro dans R . Si la conclusion de la porte $\&$ de B est étiquetée par $E \& F$, alors nous parlerons de la première (resp. la deuxième) **composante de** B pour indiquer (avec les notations de la définition) celui parmi R_i^1 et R_i^2 qui a E (resp. F) parmi ses conclusions, que l'on notera $\gamma_1(B)$ (resp. $\gamma_2(B)$). Nous écrirons aussi $\gamma_i(B)$, en sous-entendant toujours $i \in \{1, 2\}$; et nous dirons aussi de $\gamma_i(B)$ que c'est la composante i de B .

Soit LL'_2 le calcul des séquents linéaire du second ordre avec la règle mix. On montre par induction sur la preuve de LL'_2 que, modulo le renommage éventuel de certaines variables (le renommage dont il était question dans la remarque 1.2.4), à toute preuve on peut associer un réseau.

La réciproque est notoirement un résultat difficile. Cependant, le lecteur ayant une certaine familiarité avec les réseaux n'aura pas de mal à se convaincre du fait que le théorème suivant est vrai, et nous ne le priverons pas de ce plaisir...(d'autant plus que nous montrerons bientôt un "meilleur" résultat).

1.2.8. THÉORÈME. Si R est un réseau, alors il existe une preuve π de LL'_2 t.q. R est le réseau associé à π . π est alors appelée une séquentialisation de R .

1.3 Les réseaux avec sauts

La définition 1.2.7 est une définition très générale de réseau, et on pourrait essayer de prouver des résultats tels que la normalisation forte ou la confluence, relativement à cette notion de réseau. Sans additifs, ceci peut être fait en adaptant à ce cadre les démonstrations de [Girard 87] et de

[Danos 90]. Mais dès qu'on rajoute les connecteurs additifs, on est fatalement amené à définir une notion (plus ou moins générale) de *sous-réseau*, à cause de l'étape élémentaire de réduction commutative additive. La présence de graphes de correction non-connexes rend cette notion délicate à définir ; la notion naïve (un sous-réseau est une sous-structure de preuve qui est un réseau) pouvant donner lieu à des graphes incorrects : la pré-structure de preuve obtenue à partir d'un réseau en éliminant une coupure commutative additive peut ne plus être un réseau. Que le lecteur qui désire s'en convaincre simplement considère le réseau B qui est une boîte additive dont une porte coad a comme conclusion une arête étiquetée par A , et le réseau T constitué par deux liens axiome tous deux de conclusions $A, A^\perp$ et coupés entre eux. Qu'il considère ensuite le réseau R obtenu en coupant B et T . Le sous-graphe T' de T constitué des deux liens axiome (mais sans le lien coupure) est un réseau (dans le sens de la définition 1.2.7) : en le dupliquant (c.à.d. en effectuant une e.e.r. [ccad], voir la définition 2.2.3 du chapitre suivant) on aboutit à une boîte additive où les conclusions de deux portes sont prémisses d'un même lien coupure. Nous n'avons donc plus affaire à un réseau (dans le sens de la définition 1.2.7). Il faut alors trouver un moyen de définir un réseau (et donc un sous-réseau) comme un pré-réseau dont tous les graphes de correction sont, non seulement acycliques, mais aussi connexes.

Puisque la seule cause de disconnexité dans les graphes de correction est la présence des liens affaiblissement, nous introduisons la notion de **saut**, dont le but est celui de "connecter" chaque lien affaiblissement d'un réseau à une de ses boîtes (additives ou exponentielles) ou à un de ses liens axiome. L'idée d'introduire la notion de sauts pour les affaiblissements n'est pas nouvelle : que nous sachions, c'est une idée de Jean-Yves Girard, dont il est explicitement question dans l'appendice de [Girard 95a]. Cependant, les travaux faisant un usage précis de cette idée dans un cadre général tel que le notre sont très rares. A notre connaissance, seul le travail [Laurent 99] utilise explicitement une notion de saut (dans le cas restreint de la logique linéaire polarisée), principalement dans le but de prouver un théorème de séquentialisation (dans le style des théorèmes 1.2.8 et 1.3.10).

1.3.1. DÉFINITION. Soit R un réseau et B une boîte additive de R . Soient l_1 et l_2 (resp. a_1 et a_2 , B_1 et B_2) deux liens (resp. arêtes, boîtes) de R . Nous dirons que B **sépare** l_1 et l_2 (resp. a_1 et a_2 , B_1 et B_2), lorsqu'ils (elles) sont contenu(e)s dans deux composantes différentes de B . Nous dirons aussi que l_1 et l_2 (resp. a_1 et a_2 , B_1 et B_2) sont **séparé(e)s** par B , ou tout simplement **séparé(e)s**.

1.3.2. REMARQUE. Il est bien évident que si deux liens, arêtes ou boîtes d'un réseau R sont séparés, alors il existe une unique boîte (additive) qui les sépare, et nous pourrions donc à juste titre parler de *la* boîte qui les sépare.

Dans la définition suivante, la condition (1) va nous assurer qu'après une étape élémentaire de réduction quelconque (voir chapitre suivant) chaque lien affaiblissement est encore connecté à une boîte ou à un lien axiome, la condition (2) va nous garantir qu'il n'y a pas "trop" de sauts, et la condition (3) est clairement nécessaire pour le théorème de séquentialisation (théorème 1.3.10).

1.3.3. DÉFINITION. Soit T un pré-réseau, W l'ensemble des liens affaiblissement de T , Ax l'ensemble des liens axiome de T , et $\mathcal{P}(Ax)$ l'ensemble des parties de Ax . Soit j_T une fonction de domaine W et codomaine $\mathcal{P}(Ax)$.

On dit que (T, j_T) est un **pré-réseau avec sauts** lorsque pour tout lien n de W , il existe un entier $k \geq 1$ t.q. $j_T(n) = \{y_1, \dots, y_k\}$, et les conditions suivantes sont satisfaites :

(1) si $y \in j_T(n)$ est un lien contenu dans une boîte additive B qui ne contient pas n , alors il existe $y' \in j_T(n)$ t.q. B sépare y et y' .

(2) si $i, j \in \{1, \dots, k\}$ et $i \neq j$, alors il existe une boîte additive B_{ij} qui ne contient pas n et qui sépare y_i et y_j .

(3) si B est une boîte (additive ou exponentielle) et n est un lien de B , alors tout lien de $j_T(n)$ est un lien de B .

Les liens axiome de $j_T(n)$ sont appelés les **sauts** de n , et la fonction j_T est appelée fonction de saut.

1.3.4. REMARQUE. On utilise les notations de la définition précédente.

(i) Si $k \geq 2$ alors il existe une boîte additive de T contenant tous les liens de $j_T(n)$. Dans ce cas, on notera B_n la plus petite parmi ces boîtes. Il est facile de vérifier que B_n ne contient pas n .

(ii) Si y est un saut du lien affaiblissement n , alors $j_T(n) = \{y\}$ ssi n et y sont contenus dans les mêmes boîtes additives de T .

On va maintenant chercher à donner au lecteur une explication "intuitive" de la définition ci-dessus.

Soit π une preuve de LL_2 ne contenant aucune occurrence de la règle *Hyp*. On peut associer à π un (unique) réseau (dans le sens de la définition 1.2.7) R_π . Nous allons décrire comment la preuve π définit aussi une fonction de saut j_{R_π} qui fait de (R_π, j_{R_π}) un pré-réseau avec sauts (en fait, il s'agira d'un

réseau avec sauts, dans le sens cette fois de la définition 1.3.7). La description qui suit montre bien que j_{R_π} n'est pas (en général) unique.

Soit S une règle affaiblissement de π et w_S le lien de R_π associé à S . Il faut définir un ensemble $j_{R_\pi}(w_S)$ de liens axiome de R_π . Plaçons-nous dans l'arbre de preuve π sur la règle S , et imaginons de “remonter” les différentes branches de π jusqu'aux axiomes. Plus précisément : lorsqu'on rencontre une règle avec une seule prémisses on continue, lorsqu'on rencontre une règle avec deux prémisses (c.à.d. une règle $\&$ ou une règle $\otimes$), si c'est une règle $\otimes$ on choisit (arbitrairement) une des deux façons possibles de continuer à “remonter” l'arbre de preuve, tandis que si c'est une règle $\&$ on continue à “monter” en suivant les deux branches de l'arbre de preuve. Ce processus se termine en un certain nombre (forcément non nul, à cause de l'absence de règles *Hyp* dans π) de règles axiome. L'ensemble $j_{R_\pi}(w_S)$ est l'ensemble des liens axiome de R_π associés à ces règles.

Réciproquement, supposons maintenant que T ne soit pas seulement un pré-réseau avec sauts, mais aussi un réseau (suivant maintenant la définition 1.3.7 et le théorème 1.3.10).

Si $j_T(n) = \{y\}$, et y et n sont contenus dans les mêmes boîtes (additives et exponentielles) de T , alors pour chaque séquentialisation de T la règle affaiblissement correspondant à n “suivra” la règle axiome correspondant à y .

Si $j_T(n) = \{y\}$, et y est contenu dans une boîte qui ne contient pas n , alors (à cause de la condition (1) de la définition) la boîte en question est forcément exponentielle. Soit $B^!$ la plus grande parmi ces boîtes et m la pal de $B^!$. Pour toute séquentialisation de T , la règle correspondant à n “suivra” la règle bien sûr correspondant à m .

Si $k \geq 2$, alors soit m la porte $\&$ ou bien la pal de la plus grande boîte B de T ne contenant pas n et qui contient $\{y_1, \dots, y_k\}$ (B peut être la boîte additive B_n). Pour toute séquentialisation de T , la règle affaiblissement correspondant à n “suivra” la règle avec ou bien sûr correspondant à m .

Grâce à la notion de saut, on a “connecté” chaque lien affaiblissement à une boîte ou à un lien axiome : on va aboutir à une notion de réseau pour laquelle les graphes de correction ne sont plus seulement acycliques mais aussi connexes (voir définition 1.3.7). Ceci permet de définir une notion plus fine de sous-réseau (voir définition 1.3.12) qui conduit à une définition correcte de l'étape élémentaire de réduction commutative additive (voir théorème 2.3.5).

1.3.5. DÉFINITION. (graphe apparié 2) Soit (R, j_R) un pré-réseau avec

sauts. On va associer à (R, j_R) un ensemble $App(R)$ et un graphe apparié que nous noterons encore $R_{ap} = (G_R, App(R))$, et qui est obtenu à partir du graphe G_R défini en 1.2.5 en rajoutant une arête en plus pour chaque lien affaiblissement n de R de profondeur additive et exponentielle zéro :

- si $j_R(n) = \{y\}$ et y est à profondeur additive et exponentielle zéro dans R , alors nous rajoutons une arête ayant comme source n et comme but y
- autrement, nous rajoutons une arête ayant comme source n et comme but le noeud H de G_R substituant la boîte de profondeur additive et exponentielle zéro dans R qui contient tous les liens de $j_R(n)$.

1.3.6. DÉFINITION. (graphe de correction 2) La définition de switching et de graphe de correction est identique à celle donnée en 1.2.6, à la différence que R_{ap} (donc G_R) sera celui de la définition 1.3.5 et non pas celui de la définition 1.2.5.

La définition suivante est la définition définitive de réseau, à laquelle (sauf mention explicite du contraire) nous ferons toujours référence dans la suite. PN_2 est l'ensemble de tous les réseaux.

1.3.7. DÉFINITION. (réseau -définitive-) Soit (R, j_R) un pré-réseau avec sauts et $B_1, \dots, B_k$ les boîtes de profondeur additive et exponentielle zéro dans R . Nous dirons que (R, j_R) est un **réseau** lorsque les conditions suivantes sont satisfaites :

- R satisfait (ACC) : pour tout switching S de R , le graphe de correction $S(R)$ est acyclique et connexe
- pour toute boîte exponentielle $B_i \in \{B_1, \dots, B_k\}$, soit R_i le pré-réseau contenu dans B_i , et soit j_{R_i} la restriction de j_R aux liens affaiblissement de R_i . On demande que (R_i, j_{R_i}) soit un réseau
- pour toute boîte additive $B_i \in \{B_1, \dots, B_k\}$, soient R_i^1 et R_i^2 les deux pré-réseaux de la définition 1.2.7, et soit $j_{R_i^1}$ (resp. $j_{R_i^2}$) la restriction de j_R aux liens affaiblissement de R_i^1 (resp. R_i^2). On demande que $(R_i^1, j_{R_i^1})$ et $(R_i^2, j_{R_i^2})$ soient deux réseaux.

Dans la suite, nous mentionnerons rarement la fonction j_R et nous parlerons le plus souvent d'un "réseau R ", en sous-entendant toujours qu'il s'agit en fait d'un "réseau (R, j_R) ".

Comme dans le cas de LL'_2 , on peut montrer par induction que, modulo le renommage éventuel de certaines variables, à chaque preuve de LL_2 (le calcul des séquent linéaire du second ordre sans la règle mix) est associé un réseau.

1.3.8. DÉFINITION. (noeud scindant) Soit R un réseau et n un lien de R . On dira que n est scindant pour R_{ap} lorsque n satisfait une des deux conditions qui suivent :

- n est un lien cut ou $\otimes$ de conclusion a , et il existe deux sous-graphes disjoints de G_R G_1 et G_2 , t.q. G_R peut s'obtenir en connectant une arête conclusion de G_1 et une arête conclusion de G_2 par le biais du lien n (et de sa conclusion dans le cas du $\otimes$)
- n est un lien $\forall$, $\wp$ ou $?co$ et il existe deux sous-graphes disjoints de G_R G_1 et G_2 dont l'union est G_R , tels que G_1 contient n mais ne contient pas sa conclusion qui est contenue dans G_2 , et l'unique arête de G_R connectant un noeud de G_1 et un noeud de G_2 est la conclusion de n .

La proposition suivante est une conséquence directe de la proposition 4-5 p.22 de [Danos 90].

1.3.9. PROPOSITION. (n -uple scindant) Si R est un réseau contenant au moins un lien $\forall$, $\wp$ ou $?co$ de profondeur additive et exponentielle zéro, alors un parmi ces liens est scindant pour R_{ap} .

1.3.10. THÉORÈME. (séquentialisation) Si (R, j_R) est un réseau, alors il existe une preuve π de LL_2 t.q. (R, j_R) est le réseau associé à π . La preuve π est alors appelée une séquentialisation de (R, j_R) .

Démonstration : Soient les entiers p , k et q définis à partir de R par :

$p :=$ profondeur maximale des liens de R

$k :=$ nombre de liens $\forall$, $\wp$, $?co$ de R de profondeur zéro

$q :=$ nombre de liens de R de profondeur zéro.

On procède par induction sur les triplets (p, k, q) ordonnés lexicographiquement.

On va d'abord démontrer le théorème pour $k = 0$ (cas 1).

(1.1) : Si une parmi les conclusions de R est conclusion d'un lien $?w$, $\oplus$ ou $?de$, alors le résultat est une conséquence immédiate de l'hypothèse d'induction (q va diminuer).

En guise d'exemple considérons le cas de l'affaiblissement : appelons R' la structure de preuve obtenue à partir de R en effaçant le lien affaiblissement n et sa conclusion a étiquetée par $?A$. R' est un pré-réseau, et on peut considérer la restriction $j_{R'}$ de j_R à R' (on a tout simplement effacé les sauts de n). R' satisfait la condition (ACC) de la définition 1.3.7. Par hypothèse d'induction il existe donc une preuve π' de LL_2 ayant comme conclusions le séquent des étiquettes des conclusions de R' . La preuve π cherchée est alors

obtenue en appliquant au séquent conclusion de π' une règle affaiblissement de conclusion $?A$.

(1.2) : Si aucune conclusion de R n'est conclusion d'un lien $?w$, $\oplus$ ou $?de$, et si il y a dans G_R un lien $\otimes$ ou bien un lien cut , alors il est facile de voir que l'acyclicité des graphes de correction implique l'existence d'un lien $\otimes$ ou cut scindant. On applique alors comme ci-dessus l'hypothèse d'induction, et on obtiendra la preuve cherchée en appliquant aux deux preuves fournies par l'hypothèse d'induction une règle $\otimes$ ou bien une coupure, selon le cas.

(1.3) : Si aucune conclusion de R n'est conclusion d'un lien $?w$, $\oplus$ ou $?de$ et si il n'y a pas dans G_R de liens $\otimes$ ni cut mais il y a un lien axiome, alors par connexité des graphes de correction le lien axiome est le seul lien de R , et la conclusion est évidente.

(1.4) : Autrement, toutes les conclusions de R sont des conclusions de portes de boîte. Par connexité (et par la condition d'emboîtement de la définition 1.2.1), R est alors une boîte B . B contient un (resp. deux) réseau(x) si c'est une boîte exponentielle (resp. additive), au(x)quel(s) on applique l'hypothèse d'induction (cette fois c'est p qui a diminué) : on obtient une (resp. deux) preuve(s) π' (resp. π_1 et π_2) à laquelle (resp. auxquelles) on appliquera la règle d'introduction du bien sûr (resp. du $\&$).

Passons maintenant au cas vraiment significatif dans lequel $k \neq 0$ (cas 2). Par la proposition 1.3.9, il existe un lien $\wp$, $?co$ ou bien $\forall$ qui est scindant pour R_{ap} . Appelons n ce lien.

Les trois cas possibles pour n se traitent de la même façon, mais pour fixer les idées nous supposons que n est un lien $\forall$ dont la prémisse a_1 est étiquetée par $A[X]$ et la conclusion a_2 par $\forall X A$. Soient G_1 et G_2 les deux sous-graphes de G_R fournis par la définition 1.3.8, notons γ (resp. δ) les arêtes conclusions de G_1 (resp. de G_2), et Γ (resp. Δ) leurs étiquettes. Appelons G'_1 le graphe obtenu à partir de G_1 en effaçant le noeud n , et G'_2 le graphe obtenu à partir de G_2 en substituant l'arête a_2 par un lien H ayant cette même arête comme conclusion.

Puisque n est scindant, il existe une pré-structure de preuve R_1 de conclusions a_1, γ et une pré-structure de preuve R_2 de conclusions δ t.q. :

- pour tout l lien $?w$ de R , l est un lien de R_1 (et alors tous les liens de $j_R(l)$ sont des liens de R_1) ou bien l est un lien de R_2 (et alors tous les liens de $j_R(l)$ sont des liens de R_2)
- si on appelle j_{R_1} (resp. j_{R_2}) la restriction de j_R aux liens $?w$ de R_1 (resp. de R_2), alors (R_1, j_{R_1}) (resp. (R_2, j_{R_2})) est un pré-réseau avec sauts
- $R_{1ap} = (G'_1, App(R_1))$, $R_{2ap} = (G'_2, App(R_2))$ et $App(R)$ est l'union disjointe de $App(R_1)$, $App(R_2)$ et du m -uplet de toutes les arêtes de G_R incidentes en n .

La seule chose qui n'est pas immédiate à vérifier est que R_1 et R_2 satisfont la condition (i) de la définition 1.2.3 de pré-réseau : soit C l'étiquette d'une conclusion de R_1 . Si C n'est pas l'étiquette $A[X]$ de a_1 , alors aucune variable libre de C n'est la variable propre d'un lien $\forall$ de R (donc à fortiori de R_1). Si C est l'étiquette $A[X]$ de a_1 , alors soit Y une variable libre de C . Si $Y \neq X$ et si Y est la variable propre d'un lien $\forall$ de R_1 , alors la conclusion a_2 du lien n ne peut pas être conclusion de R (par (i) de la définition de pré-réseau). Soit alors m le lien de R dont a_2 est une prémisse. On aurait alors, dans R_{ap} , une arête entre le lien $\forall$ de R_1 de variable propre Y et m , ce qui contredit le fait que n est scindant.

Quant au cas $Y = X$, si X était la variable propre d'un lien $\forall$ de R_1 , alors ce serait la variable propre de deux liens $\forall$ différents de R , ce qui contredit (ii) de la définition 1.2.3. Le cas des conclusions de R_2 est immédiat.

R_1 et R_2 satisfont la condition (ACC) de la définition 1.3.7. Par hypothèse d'induction il existe donc une preuve π_1 (resp. π_2) de LL_2 ayant comme conclusions le séquent $\vdash \Gamma, A[X]$ (resp. $\vdash \Delta$). La preuve π cherchée est alors obtenue en appliquant au séquent conclusion de π_1 une règle d'introduction du $\forall$ de conclusion $\vdash \Gamma, \forall X A$, et en substituant la preuve ainsi obtenue à la règle *Hyp* de conclusion $\forall X A$ dans π_2 . Pour se convaincre que π est bien une preuve de LL_2 , il suffira de prouver :

- (a) il n'y a pas d'occurrence de X libre dans Γ
 - (b) toute règle d'introduction du $\forall$ dans π_2 n'utilise pas comme variable propre une variable apparaissant libre dans Γ ou dans $\forall X A$.
- (a) et (b) sont toutes les deux des conséquences immédiates du fait que R est un pré-réseau, et du fait n est scindant. $\square$

1.3.11. REMARQUE. Il est bien évident que tout réseau associé à une preuve de LL_2 ne faisant pas usage de la règle *Hyp* est un réseau sans liens H , la réciproque étant tout aussi évidente. C'est donc à cette notion de réseau que nous ferons référence dans la suite de la thèse.

1.3.12. DÉFINITION. Soit (T, j_T) un réseau et R un sous-graphe de T . Nous dirons que R est un **sous-réseau** de T , quand (R, j_R) est un réseau, où j_R est la restriction de j_T à l'ensemble des liens affaiblissement de R .

1.3.13. REMARQUE. (i) Si R est un sous-réseau du réseau T et n est un lien affaiblissement de R , alors les liens de $j_R(n) = j_T(n)$ sont des liens de R .

(ii) Si B est une boîte additive du réseau (R, j_R) , alors (voir définition 1.3.7) $(\gamma_i(B), j_{\gamma_i(B)})$ (resp. (B, j_B)), où $j_{\gamma_i(B)}$ (resp. j_B) est la restriction à $\gamma_i(B)$ (resp. B) de j_R , est un sous-réseau de (R, j_R) .

Le lemme suivant est une conséquence des définitions 1.3.7 et 1.3.12 de réseau et sous-réseau, et nous en ferons usage dans l'annexe A. Le lecteur ne manquera certainement pas de remarquer que ce lemme n'est pas vrai si l'on fait référence à la notion préliminaire de réseau fournie par la définition 1.2.7.

1.3.14. LEMME. Soit (T, j_T) un réseau. Appelons T_j le graphe obtenu à partir de T en dessinant les arêtes connectant les liens affaiblissement de T avec leurs sauts. Si (R, j_R) est un sous-réseau de (T, j_T) , et si m et n sont deux liens de R , alors il existe un chemin (orienté) de T_j ayant m comme lien initial et n comme lien terminal, et qui ne traverse que des liens de R .

Chapitre 2

Normalisation des réseaux de PN_2

Nous allons nous occuper, dans ce chapitre, des étapes élémentaires de réduction (e.e.r.) associées aux liens coupure de PN_2 . Celles-ci seront des déformations des graphes que sont les réseaux. Une fois la procédure d'élimination des coupures définie, il s'agira de montrer que nos étapes de normalisation préservent la correction (théorème 2.3.5).

Le cadre très général dans lequel nous nous sommes placés nous empêche d'invoquer des théorèmes similaires déjà démontrés.

2.1 La notion de substitution

2.1.1. DÉFINITION. (variante de [DanReg 95]) Soit R une préstructure de preuve. Un chemin de R est une suite d'arêtes ou d'arêtes renversées (c.à.d. qu'un chemin peut parcourir une arête de son but vers sa source). Nous noterons $\alpha, \beta, \dots$ les arêtes d'une préstructure de preuve (orientées comme d'habitude suivant la définition 1.2.1) et $\alpha^*, \beta^*, \dots$ les arêtes précédentes "renversées" (c.à.d. orientées maintenant dans la direction opposée). Soit a (resp. b) une arête ou une arête renversée dont le but (resp. la source) est le lien n ; nous écrirons ab pour indiquer le chemin constitué de l'arête a suivie du lien n , lui-même suivi de l'arête b .

Nous dirons que le chemin Φ de R est un **chemin droit** si :

- (i) Φ ne contient pas de sous-chemins $\alpha^*\alpha$ ni $\alpha\alpha^*$
- (ii) Si α et β sont deux prémisses d'un même lien n et si $\alpha\beta^*$ est un sous-chemin de Φ , alors n est un lien coupure

(iii) Φ ne contient pas de sous-chemins $\alpha^*\beta$, où α et β sont deux conclusions d'un même lien H .

Dans la suite de la thèse, sauf mention explicite du contraire, lorsque nous écrirons “chemin” nous ferons toujours référence à la notion de “chemin droit” de la définition précédente.

Nous parlerons souvent d'un “chemin ayant n comme premier lien et m comme lien terminal” : il s'agira du chemin dont la première arête est celle qui vient “après” n (relativement à l'orientation du chemin en question) et dont la dernière arête est celle qui vient “avant” m (toujours relativement à l'orientation du chemin en question). Nous parlerons aussi d'un “chemin ayant n comme premier lien et a comme arête terminale” : il s'agira du chemin dont la première arête est celle qui vient “après” n et dont la dernière arête est a . Si Φ et Ψ sont deux chemins d'un réseau R , nous noterons $\Phi \cap \Psi$ l'ensemble des liens et des arêtes (orientées) qui appartiennent à Φ et à Ψ .

2.1.2. DÉFINITION. Soit R un préréseau, et soit Φ un chemin de R ayant n comme premier lien et m comme lien terminal (resp. a comme arête terminale). Nous dirons que Φ est un **chemin direct** ssi les deux conditions suivantes sont satisfaites :

- (i) Φ ne contient aucun lien *cut* (excepté éventuellement m), et aucun lien axiome (excepté éventuellement n)
- (ii) chaque arête de Φ est orienté du haut vers le bas (i.e. il s'agit d'une arête -et pas d'une arête renversée- de R).

Intuitivement, un chemin change de direction lorsqu'il traverse un lien axiome ou coupure, et il peut contenir des arêtes descendantes et des arêtes ascendantes. Ce n'est pas le cas d'un chemin direct : il ne peut pas changer de direction (condition (i)) et toutes ses arêtes sont descendantes (condition (ii)). On peut dire (de façon un peu imprécise) que si A est l'étiquette d'une conclusion d'un réseau ou bien de la prémisse d'un lien coupure, alors un chemin direct est un chemin issu d'une sous-formule de A et ayant A comme arête terminale.

2.1.3. REMARQUE. Soit R un réseau et m un lien de R . Une et une seule des deux conditions suivantes est satisfaite :

- (1) il existe un lien coupure c de R et un (unique) chemin direct Φ_m ayant m comme premier lien et c comme lien terminal
- (2) il existe une arête conclusion a de R et un (unique) chemin direct Φ_m ayant m comme premier lien et a comme arête terminale.

Intuitivement, tout lien d'un réseau R "est issu" d'un lien logique, axiome ou affaiblissement, comme l'exprime plus précisément le lemme qui suit. Nous dirons alors que le lien n du réseau R est un **lien principal** de R si n est un lien logique, axiome ou affaiblissement.

2.1.4. LEMME. Soit T un réseau. Pour tout lien n de T , il existe un lien principal l et un chemin direct Φ de T ayant l comme premier lien et n comme lien terminal.

Démonstration : Evident. □

Nous utiliserons le lemme suivant pour définir les étapes élémentaires de réduction des réseaux, ainsi que la notion de substitution d'un sous-réseau à un autre dans un réseau.

2.1.5. LEMME. Soit B une boîte additive du réseau (R, j_R) et soit t une porte de B .

Il existe un entier $h \geq 2$ et un ensemble $Ax_t = \{z_1, \dots, z_h\}$ de liens axiome de B t.q. les conditions suivantes sont satisfaites :

- (i) si $z \in Ax_t$ est un lien d'une boîte additive B' contenue dans B (c'est en particulier le cas pour B), il existe $z' \in Ax_t$ t.q. B' sépare z' et z
- (ii) si $z_i, z_j \in \{z_1, \dots, z_h\}$ et $z_i \neq z_j$, il existe une boîte additive B_{ij} contenue dans B (on peut avoir $B_{ij} = B$) qui sépare z_i et z_j
- (iii) $\forall i \in \{1, \dots, h\}$ il existe un chemin direct Ψ_t ayant t comme lien terminal t.q. z_i satisfait une des deux conditions suivantes :
 - (α) Ψ_t a z_i comme lien initial
 - (β) il existe un lien affaiblissement m_i de B t.q. $z_i \in j_R(m_i)$ et Ψ_t a m_i comme lien initial.

Démonstration : Soit L l'ensemble des liens axiome ou affaiblissement l de B t.q. il existe un (unique) chemin direct Φ_l de T ayant l comme premier lien et t comme lien terminal. On note $\sharp\Phi_l$ le nombre de portes de boîtes additives (t comprise) traversées par Φ_l . Evidemment $L \neq \emptyset$.

On procède par induction sur $p = \max\{\sharp\Phi_l : l \in L\}$.

Si $p = 1$, alors pour $i \in \{1, 2\}$ il existe $y_i \in \gamma_i(B)$ t.q. l'unique porte de boîte additive traversée par Φ_{y_i} est t (en général il y aura bien sûr plusieurs de ces liens et nous choisissons un d'entre eux pour chacune des composantes de B). L'ensemble $Ax_t = Y_1 \cup Y_2$ fera l'affaire, où pour $i \in \{1, 2\}$ on a $Y_i = \{y_i\}$ si y_i est un lien axiome, et $Y_i = j_R(y_i)$ si y_i est un lien affaiblissement. Remarquons que $\Psi_t = \Phi_{y_i}$, et si y_i est un lien axiome, alors y_i satisfait la

condition (α) , tandis que si y_i est un lien affaiblissement alors chaque lien $z \in j_R(y_i)$ satisfait la condition (β) .

Si $p > 1$ il existe un lien $u \in L$ t.q. $\# \Phi_u > 1$. Supposons par exemple que $u \in \gamma_1(B)$. Soit B_1 la plus grande boîte additive de $\gamma_1(B)$ contenant u et soit t_1 la porte de B_1 traversée par Φ_u . Soit L_1 l'ensemble des liens axiome ou affaiblissement l de B_1 t.q. il existe un (unique) chemin direct Φ_l^1 de T ayant l comme premier lien et t_1 comme lien terminal. Tout lien l de L_1 est un lien de L , et dans ce cas Φ_l^1 est un sous-chemin de Φ_l et $\# \Phi_l^1 < \# \Phi_l$. On a $p_1 = \max\{\# \Phi_l^1 : l \in L_1\} < p$ et on peut appliquer l'hypothèse d'induction à B_1 : il existe un entier $h_1 \geq 2$ et un ensemble $Ax_{t_1} = \{z_1^1, \dots, z_{h_1}^1\}$ de liens axiome de B_1 qui satisfont les conditions du lemme.

Si il y a aussi un lien $u' \in \gamma_2(B)$ t.q. $\# \Phi_{u'} > 1$, alors appelons B_2 la plus grande boîte additive de $\gamma_2(B)$ contenant u' , t_2 la porte de B_2 traversée par $\Phi_{u'}$ et appliquons (comme avant) l'hypothèse d'induction à B_2 : il existe un entier $h_2 \geq 2$ et un ensemble $Y_2 = Ax_{t_2} = \{z_1^2, \dots, z_{h_2}^2\}$ de liens axiome de B_2 qui satisfont les conditions du lemme. Si un tel lien u' n'existe pas, alors il existe $y_2 \in \gamma_2(B) \cap L$ t.q. la seule porte de boîte additive traversée par Φ_{y_2} est t . Dans ce cas, si y_2 est un lien axiome, on pose $Y_2 = \{y_2\}$, tandis que si y_2 est un lien affaiblissement on pose $Y_2 = j_R(y_2)$. Nous allons montrer que, dans tous les cas, l'ensemble $Ax_t = Ax_{t_1} \cup Y_2$ satisfait les conditions du lemme.

Condition (i) : elle est satisfaite par hypothèse d'induction par les boîtes additives contenues dans B_1 (B_1 comprise) et donc par toute boîte additive de $\gamma_1(B)$. Si il existe $u' \in \gamma_2(B)$ t.q. $\# \Phi_{u'} > 1$, alors la condition (i) est aussi satisfaite par les boîtes additives contenues dans B_2 (B_2 comprise) et donc par toute boîte additive de $\gamma_2(B)$. Si un tel lien axiome u' n'existe pas, alors y_2 est à profondeur additive zéro dans $\gamma_2(B)$, et soit $Y_2 = \{y_2\}$ (auquel cas (i) est trivialement satisfaite par toute boîte additive de $\gamma_2(B)$), soit $Y_2 = j_R(y_2)$ et alors la condition (1) de la définition 1.3.3 permet d'affirmer que la condition (i) est satisfaite par toute boîte additive contenant un lien de Y_2 (et donc par toute boîte additive de $\gamma_2(B)$). Le fait que la boîte additive B satisfasse la condition (i) est une conséquence immédiate de la définition de Ax_t .

Condition (ii) : soient $z, z' \in Ax_t$. Si $z, z' \in Ax_{t_1}$ ou si $z, z' \in Y_2$, alors (ii) est satisfaite par hypothèse d'induction (et grâce à la condition (2) de la définition 1.3.3). Autrement, supposons que $z \in Ax_{t_1}$ et que $z' \in Y_2$: B est alors la boîte additive qui sépare z et z' .

Condition (iii) : rappelons d'abord que Φ_u a $u \in L$ comme premier lien, la porte t de B comme lien terminal, et que Φ_u traverse t_1 . Soit alors $\Phi_{t_1 t}$ le sous-chemin de Φ_u ayant t_1 comme premier lien et t comme lien terminal.

Si $z \in Ax_{t_1}$, alors par hypothèse d'induction il existe un chemin direct Ψ_{t_1} ayant t_1 comme lien terminal et t.q. Ψ_{t_1} et z satisfont (α) ou (β) . Le chemin direct Ψ_t obtenu en connectant Ψ_{t_1} et $\Phi_{t_1 t}$ a t comme lien terminal et Ψ_t et z satisfont (α) ou (β) . Si $z \in Y_2 = Ax_{t_2}$, alors on procède exactement comme pour $z \in Ax_{t_1}$. Si $z \in Y_2 = \{y_2\}$, alors y_2 est un lien axiome et il est bien clair que y_2 et Φ_{y_2} satisfont (α) . Si $z \in Y_2 = j_R(y_2)$, alors y_2 est un lien affaiblissement et tout lien de $j_R(y_2)$ et Φ_{y_2} satisfont (β) . $\square$

Soit α un sous-réseau du réseau T , et soit β un réseau avec les mêmes conclusions que α . Soit G le graphe obtenu à partir de T en substituant β à α . En changeant (éventuellement) le nom de certaines variables de G , on obtiendra un graphe satisfaisant les conditions des définitions 1.2.1 et 1.2.3. Appelons R ce préréseau.

2.1.6. DÉFINITION. Soit α un sous-réseau du réseau T , et soit β un réseau avec les mêmes conclusions que α . On peut associer au préréseau R défini ci-dessus un réseau en définissant la fonction j_R . En général, ceci peut se faire de plusieurs façons différentes. Nous noterons $T[\beta/\alpha]$ un quelconque de ces réseaux dont la fonction de saut $j_{T[\beta/\alpha]}$ satisfait les trois conditions suivantes. Soit n un lien de R . Si n est un lien de α , alors (cf définition 1.3.12) tous les liens de $j_T(n)$ sont des liens de α . On demande que :

- (1) si n est un lien de β , alors $j_{T[\beta/\alpha]}(n) = j_\beta(n)$
- (2) si n est un lien de T qui n'est pas un lien de α et si tout lien de $j_T(n)$ n'est pas un lien de α , alors $j_{T[\beta/\alpha]}(n) = j_T(n)$
- (3) si n est un lien de T qui n'est pas un lien de α , soit $j_T(n) = \{y_1, \dots, y_k\}$ et supposons que $y_1, \dots, y_l$ ($1 \leq l \leq k$) sont des liens de α tandis que $y_{l+1}, \dots, y_k$ ne sont pas des liens de α . Nous exigeons que une (et une seule) des deux conditions suivantes soit satisfaite :
 - il existe un lien axiome y de β de profondeur additive zéro dans β et $j_{T[\beta/\alpha]}(n) = \{y\} \cup \{y_{l+1}, \dots, y_k\}$
 - il existe une boîte additive B de profondeur additive zéro dans β et un ensemble $Ax = \{z_1, \dots, z_h\}$ de liens axiome de B satisfaisant les conditions (i) et (ii) du lemme 2.1.5, et on a $j_{T[\beta/\alpha]}(n) = Ax \cup \{y_{l+1}, \dots, y_k\}$.

2.1.7. REMARQUE. (i) On peut vérifier que tout préréseau $T[\beta/\alpha]$ défini ci-dessus est un préréseau avec sauts qui satisfait (inductivement) la condition (ACC) de la définition 1.3.7, c.à.d. que c'est un réseau.

(ii) Pour tout T , α , et β satisfaisant les hypothèses de la définition ci-dessus, il existe (en vertu du lemme 2.1.5) au moins un réseau $T[\beta/\alpha]$.

(iii) Soient T , α , et β comme dans la définition précédente. Soit δ un sous-réseau de T qui n'est pas un sous-réseau de α , et soit S_δ la préstructure de preuve obtenue à partir de δ en substituant β à α (si α n'est pas un sous-réseau de δ , alors $S_\delta = \delta$). Pour tout réseau $T[\beta/\alpha]$, on note $\delta[\beta/\alpha]$ le réseau associé à S_δ t.q. la fonction $j_{\delta[\beta/\alpha]}$ soit la restriction de la fonction $j_{T[\beta/\alpha]}$ aux liens affaiblissements de S_δ : $\delta[\beta/\alpha]$ est alors un sous-réseau de $T[\beta/\alpha]$.

2.2 La normalisation

Pour tout lien coupure x du réseau R , nous allons d'abord définir le graphe R' obtenu à partir de R en éliminant le lien coupure x (définition 2.2.3), et nous définirons ensuite la fonction $j_{R'}$ obtenue à partir de j_R en éliminant le lien coupure x (définition 2.2.8). Après quoi, il faudra se convaincre que $(R', j_{R'})$ est un réseau (théorème 2.3.5).

2.2.1. DÉFINITION. Soit R un réseau et x un lien coupure de R .

x est appelé un lien coupure **logique** quand les deux prémisses du lien sont conclusions de deux liens logiques (nécessairement duaux) : on écrit aussi que ces liens coupure sont des liens $\wp / \otimes$, $!/?de$, $\&/\oplus_1$, $\&/\oplus_2$ (ou simplement $\&/\oplus$), $\forall/\exists$.

x est appelé un lien coupure **contraction** (resp. **affaiblissement**) lorsque une prémisses de x est conclusion d'un lien contraction (resp. affaiblissement) et l'autre est conclusion d'un lien bien sûr : on note aussi co (resp. w) un tel lien coupure.

x est appelé un lien coupure **commutative exponentielle** lorsque une prémisses de x est conclusion d'un lien pax et l'autre est conclusion d'un lien bien sûr : on notera aussi cc un tel lien coupure.

x est appelé un lien coupure **axiome** lorsque une des deux prémisses de x est conclusion d'un lien axiome.

x est appelé un lien coupure **commutative additive** lorsque une des deux prémisses de x est conclusion d'un lien coad : on notera aussi $ccad$ un tel lien coupure.

2.2.2. REMARQUE. Si x est un lien coupure d'un réseau, alors x est forcément d'un seul type parmi ceux définis ci-dessus, *sauf dans un cas* : si une des prémisses de x est conclusion d'un lien axiome et l'autre est conclusion d'un lien coad, alors x est à la fois un lien coupure axiome et un lien coupure $ccad$.

On va maintenant définir les étapes élémentaires de réduction (e.e.r.) associées aux différents types de liens coupure. L'application successive de plusieurs e.e.r. est ce qu'on appelle le processus **d'élimination des coupures** ou de **normalisation**. Si après avoir appliqué au réseau R un certain nombre d'e.e.r. on aboutit à un réseau R_0 qui ne contient pas de liens coupures, nous dirons que R_0 est une **forme normale** de R .

2.2.3. DÉFINITION. Soit R un réseau et x un lien coupure de R .

Si x est un lien coupure logique (mais pas $\forall/\exists$), contraction, ou *cc*, alors on note $[x]$ l'unique e.e.r. associée à x et définie dans [Girard 87]. Lorsque x est un lien coupure logique, on appelle $[x]$ une **e.e.r. logique**.

Si x est un lien coupure $\forall/\exists$ dont une prémisse est étiquetée par $\forall XA$ (conclusion du lien $\forall$ de prémisse $A[X]$) et l'autre par $\exists XA^\perp$ (conclusion du lien $\exists$ de prémisse $A^\perp[B/X]$), alors l'e.e.r. $[\forall/\exists]$ consiste d'abord à substituer *toutes* les occurrences libres de X dans R par la formule B , et ensuite le lien coupure entre les arêtes étiquetées par $\forall XA$ et $\exists XA^\perp$ par un lien coupure entre $A[B/X]$ et $A^\perp[B/X]$ (cette étape n'est autre que celle de [Girard 91a] adaptée au système PN_2 tout entier).

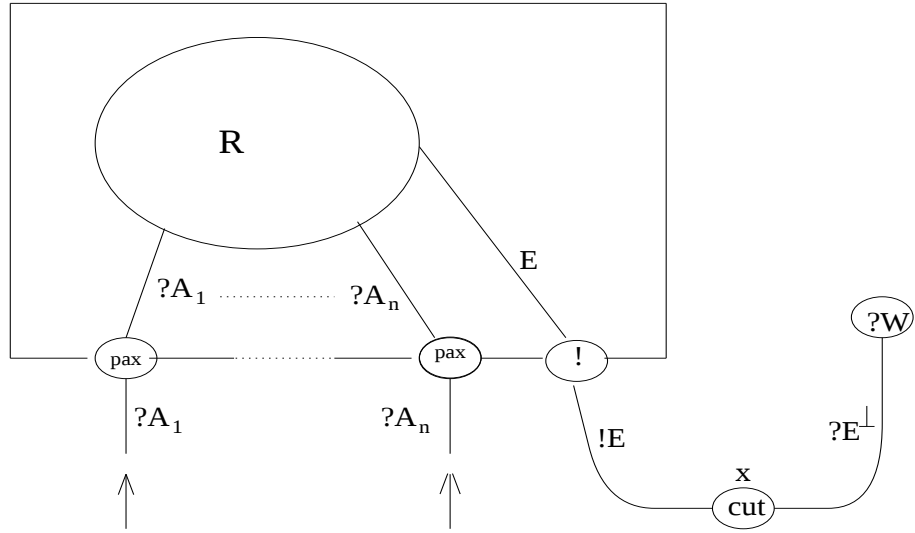
Si x est un lien coupure affaiblissement, alors on note $[x]$ l'unique e.e.r. associée à x et définie dans la figure 3.

Si x est un lien coupure axiome, alors on note $[ax]$ l'e.e.r. axiome associée à x et définie dans [Girard 87]. (Lorsque les deux prémisses de x sont conclusions de liens axiome, on a en fait *deux* e.e.r. possibles).

Si x est un lien coupure *ccad*, soit n un des liens *coad* dont la conclusion est prémisse de x (il y a au moins un tel lien et au plus deux tels liens) et appelons A l'étiquette de la conclusion de n . Soit S un sous-réseau *quelconque* de R ayant la prémisse de x étiquetée par $A^\perp$ parmi ses conclusions, et soit B la boîte additive de R ayant la prémisse de x étiquetée par A parmi ses conclusions. On note $[ccad]$ l'e.e.r. définie dans la figure 4. Cette e.e.r. est dite *B-attractive* (ou attractive pour la boîte B) parce que c'est B la boîte additive de R qui avale (et duplique) le sous-réseau S .

Dans la suite, $[x]$ indiquera une quelconque e.e.r. associée au lien coupure x .

La configuration

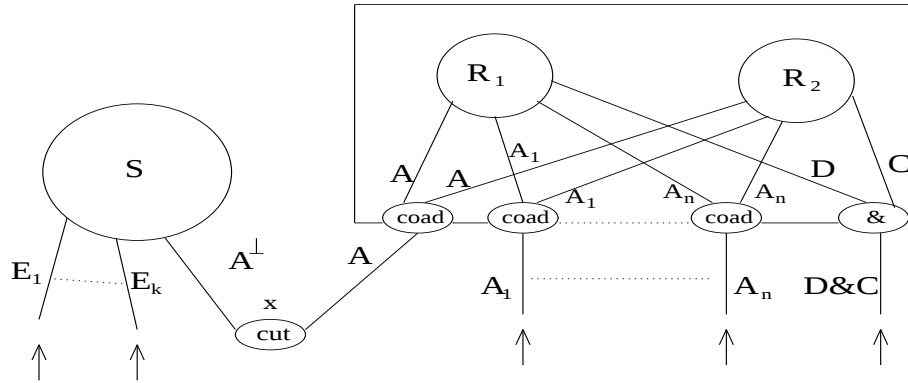


est remplacée par la configuration



Figure 3. L'étape élémentaire de réduction $[w]$.

La configuration



est remplacée par la configuration

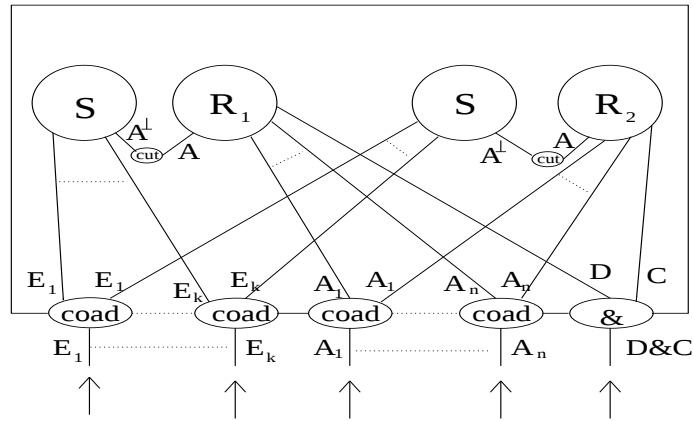


Figure 4. L'étape élémentaire de réduction $[ccad]$.

2.2.4. DÉFINITION. Soit R un réseau. On dit que R' (qui pour l'instant est simplement un graphe) est un réduit en une étape de R , si il existe un lien coupure x de R t.q. R' est obtenu à partir de R en appliquant l'e.e.r. $[x]$: on écrit dans ce cas $R \xrightarrow{[x]} R'$.

2.2.5. REMARQUE. Soit x un lien coupure du réseau R , et supposons que $R \xrightarrow{[x]} R'$. Si $[x] \neq [ccad]$, alors le graphe R' est univoquement déterminé par le lien coupure x , tandis que si $[x] = [ccad]$, pour déterminer le graphe R' encore faut-il savoir quelle est la boîte B t.q. $[x]$ est B -attractive, et quel est le sous-réseau S de R avalé par B lors de l'e.e.r. $[x]$.

Remarquons aussi que la seule e.e.r. faisant intervenir la fonction de saut j_R est l'e.e.r. $[ccad]$.

2.2.6. DÉFINITION. Soit R' le réduit en une étape du réseau R .

Tout lien axiome ou logique n de R' provient d'une unique occurrence du "même" lien axiome ou logique $\overleftarrow{n}$ de R . On appelle $\overleftarrow{n}$ l'**ancêtre du lien** n de R .

Réciproquement, on définit l'ensemble des **résidus du lien logique (resp. axiome) l de R** , comme l'ensemble de toutes les occurrences l' de liens logiques (resp. axiome) de R' telles que $\overleftarrow{l'} = l$.

Dans le cas particulier où le lien logique de R' considéré est un lien avec ou un lien bien sûr, il s'agit en fait d'une boîte additive ou exponentielle B , qui provient d'une unique occurrence de "la même" boîte additive ou exponentielle $\overleftarrow{B}$ de R (remarquons que, cependant, les liens contenus dans la boîte ne sont pas forcément les mêmes dans R et dans R'). On appelle $\overleftarrow{B}$ l'**ancêtre de la boîte** B de R .

Réciproquement, on définit l'ensemble des **résidus de la boîte B_1 de R** comme l'ensemble de toutes les occurrences B'_1 de boîtes de R' telles que $\overleftarrow{B'_1} = B_1$.

De la même façon, toute occurrence de lien coupure c de R' provient d'une unique occurrence $\overleftarrow{c}$ de lien coupure avec les mêmes formules actives (à des substitutions près si la coupure réduite est une coupure $\forall/\exists$) de R , sauf pour les liens coupure créés par une e.e.r. logique. Dans ce dernier cas, nous dirons que les liens créés n'ont pas d'ancêtres dans R et que le lien coupure réduit n'a pas de résidus dans R' .

On peut aussi définir les notions d'ancêtre et de résidu pour les liens affaiblissement.

2.2.7. DÉFINITION. Soit R un réseau et $R \xrightarrow{[x]} R'$.

Si $[x] \neq [w]$, alors tout lien affaiblissement n de R' provient d'une unique occurrence du "même" lien affaiblissement $\overleftarrow{n}$ de R , que nous appellerons (comme d'habitude) **l'ancêtre du lien** n de R . Réciproquement, on définit l'ensemble des **résidus du lien** affaiblissement l de R , comme l'ensemble de toutes les occurrences l' de liens affaiblissement de R' telles que $\overleftarrow{l'} = l$.

Si $[x] = [w]$, alors tout lien affaiblissement n de R' provient d'une unique occurrence du "même" lien affaiblissement $\overleftarrow{n}$ de R , sauf les k liens affaiblissements $n_1, \dots, n_k$ qui sont créés par l'e.e.r. $[w]$ ($k \geq 0$). Dans ce cas, si m est le lien affaiblissement de R dont la conclusion est prémisses de x , on dira que m est l'ancêtre de n_i ($i \in \{1, \dots, k\}$) et que $\{n_1, \dots, n_k\}$ est l'ensemble des résidus de m dans R' .

Soit R' le réduit en une étape de R . La définition suivante associe à R' plusieurs fonctions $j_{R'}$.

2.2.8. DÉFINITION. Soit (R, j_R) un réseau et supposons que $R \xrightarrow{[x]} R'$. On distingue tous les cas possibles pour $[x]$ et on définit $j_{R'}^{[x]}$ (qui sera simplement notée $j_{R'}$).

Soit m' un lien affaiblissement de R' , m son ancêtre dans R et soit $j_R(m) = \{y_1, \dots, y_k\}$ l'ensemble des sauts de m . On veut définir $j_{R'}(m')$.

Cas 1 : $[x]$ est une e.e.r. axiome. Soit n le lien axiome effacé par $[x]$. $\forall i \in \{1, \dots, k\}$, si $y_i \neq n$ alors y_i a un unique résidu y'_i dans T' .

Si $\forall i \in \{1, \dots, k\}$ $y_i \neq n$, alors on définit $j_{R'}(m') = \{y'_1, \dots, y'_k\}$. Autrement supposons que $n = y_1$. Soit B_1 la plus petite boîte additive contenant x et $n = y_1$ si elle existe. On appelle α la composante de B_1 contenant x et y_1 si B_1 existe, et on appelle α le réseau R autrement. y_1 est un lien de profondeur additive zéro dans α , et donc (en vertu de la condition (2) de la définition 1.3.3) $\forall i \in \{2, \dots, k\}$, y_i n'est pas un lien de α .

Il existe un lien axiome ou affaiblissement l et un chemin direct Φ_l de R ayant l comme premier lien et x comme lien terminal, qui ne traverse pas y_1 . (Remarquons que l et Φ_l ne sont pas univoquement déterminés). Soit l est contenu dans une boîte additive de R qui ne contient pas x (premier cas), soit l est contenu dans toutes les boîtes additives de R contenant x (deuxième cas).

Dans le deuxième cas, si l est un lien axiome alors on appelle l' son unique résidu dans R' et on définit $j_{R'}(m') = \{l', y'_2, \dots, y'_k\}$. Si l est un lien affaiblissement, alors tout saut de l est un lien de α , différent de y_1 (autrement un graphe de correction de R serait cyclique), qui a un unique

résidu dans R' . On appelle Z' l'ensemble de tous ces résidus et on définit $j_{R'}(m') = Z' \cup \{y'_2, \dots, y'_k\}$.

Dans le premier cas, soit B la plus grande parmi ces boîtes et soit t la porte de B traversée par Φ_t . Soit $Ax_t = \{z_1, \dots, z_h\}$ ($h \geq 2$) l'ensemble des liens axiome de R fournis par le lemme 2.1.5. $\forall j \in \{1, \dots, h\}$ z_j a un unique résidu z'_j dans R' . Soit Ax'_t l'ensemble des h résidus des liens de Ax_t . On définit $j_{R'}(m') = Ax'_t \cup \{y'_2, \dots, y'_k\}$.

Cas 2 : x est un lien coupure $\wp / \otimes$, $\forall / \exists$ ou $! / ?de$. Alors $\forall i \in \{1, \dots, k\}$, y_i a un unique résidu y'_i et on définit $j_{R'}(m') = \{y'_1, \dots, y'_k\}$.

Cas 3 : x est un lien coupure $\& / \oplus$. Alors on définit $j_{R'}(m')$ comme l'ensemble des résidus des liens de $j_R(m)$ (dans ce cas certains liens de $j_R(m)$ pourraient n'avoir aucun résidu dans R'). La condition (1) de la définition 1.3.3 est ici cruciale : elle implique que $j_{R'}(m') \neq \emptyset$.

Cas 4 : x est un lien coupure affaiblissement. Soit alors n le lien bien sûr dont la conclusion est une prémisse de x , soit $B^!$ la boîte exponentielle associée à n et soit z le lien affaiblissement dont la conclusion est une prémisse de x . Aucun lien de $j_R(z)$ n'est contenu dans $B^!$ (en vertu de la définition 1.3.7), et tout lien de $j_R(z)$ a donc un unique résidu dans R' . Soit Z' l'ensemble de ces résidus. Il est important d'insister sur le fait que dans tous les cas $Z' \neq \emptyset$. Soit J' l'ensemble des résidus des liens de $j_R(m)$. On distingue 3 cas :

(4.1) $m \neq z$ et aucun saut de m n'est contenu dans $B^!$: dans ce cas $j_{R'}(m')$ est simplement l'ensemble des résidus des éléments de $j_R(m)$

(4.2) $m \neq z$ et certains sauts de m sont contenus dans $B^!$: on pose alors $j_{R'}(m') = J' \cup Z'$

(4.3) $m = z$: dans ce cas m' est un des liens affaiblissement créés par $[x]$, et on pose $j_{R'}(m') = Z'$.

Cas 5 : $[x]$ est une e.e.r. contraction. Soit alors n le lien bien sûr dont la conclusion est une prémisse de x , soit $B^!$ la boîte exponentielle associée à n et soit $B_1^!$ et $B_2^!$ ses deux résidus dans R' . Soit Y le sous-ensemble de $j_R(m)$ contenant tous les liens de $j_R(m)$ qui sont contenus dans $B^!$, et soit Y'_1 (resp. Y'_2) l'ensemble des résidus des liens de Y qui sont contenus dans $B_1^!$ (resp. $B_2^!$). Soit Z' l'ensemble des résidus dans R' des liens de $j_R(m)$ qui ne sont pas contenus dans Y .

Si m est un lien de $B^!$, alors ($Y = j_R(m)$ et) il existe $i \in \{1, 2\}$ t.q. m' est un lien de $B_i^!$. On définit alors $j_{R'}(m') = Y'_i$.

Si m n'est pas un lien de $B^!$, alors nous distinguons le cas $Y = \emptyset$ du cas $Y \neq \emptyset$. Si $Y = \emptyset$, alors on définit $j_{R'}(m') = Z' = \{y'_1, \dots, y'_k\}$. Si $Y \neq \emptyset$, alors on définit $j_{R'}(m') = Z' \cup Y'_i$ où on a choisi l'entier $i \in \{1, 2\}$ arbitrairement.

Cas 6 : x est un lien coupure cc . Dans ce cas $\forall i \in \{1, \dots, k\}$ y_i a un unique

résidu y'_i , et on définit simplement $j_{R'}(m') = \{y'_1, \dots, y'_k\}$.

Cas 7 : $[x]$ est une e.e.r. $[ccad]$. Soit B la boîte additive de R t.q. $[x]$ est B -attractive et soit (S, j_S) le sous-réseau de (R, j_R) dupliqué par $[x]$. Soit B' le résidu de B dans R' . Soit Y le sous-ensemble de $j_R(m)$ contenant les liens de $j_R(m)$ qui sont des liens de S , et soit Y'_1 (resp. Y'_2) l'ensemble des résidus des liens de Y qui sont contenus dans la première (resp. la deuxième) composante de B' .

Si m est un lien de S , alors (en vertu de la définition 1.3.12) $j_R(m) = Y$ et il existe $i \in \{1, 2\}$ t.q. m' est un lien de $\gamma_i(B')$. On définit alors $j_{R'}(m') = Y'_i$. Si m n'est pas un lien de S , alors on définit simplement $j_{R'}(m')$ comme l'ensemble de tous les résidus des liens de $j_R(m)$.

2.2.9. REMARQUE. Soit $R \xrightarrow{[x]} R'$, où (R, j_R) est un réseau. Soit m' un lien affaiblissement de R' et m son ancêtre dans R . Alors :

- (i) Si $[x]$ n'est pas une e.e.r. axiome ni une e.e.r. affaiblissement, alors $j_{R'}(m')$ ne contient que des résidus dans R' de liens de $j_R(m)$.
- (ii) Si $[x]$ n'est pas une e.e.r. $[co]$, alors $j_{R'}(m')$ contient tous les résidus de $j_R(m)$.

2.2.10. REMARQUE. Soient R et R' deux réseaux t.q. $R \xrightarrow{[x]} R'$, et soit n un lien axiome, logique, affaiblissement, ou coupure de R . Si une des deux conditions suivantes est satisfaite, alors n n'a pas de résidu dans R' et nous dirons que n appartient au sous-réseau de R effacé par $[x]$:

- (a) $[x]$ est une e.e.r. $[w]$ et n est un lien de la boîte exponentielle effacée par $[x]$
- (b) x est un lien coupure $\&/\oplus$ dont une prémisse est la porte $\&$ de la boîte additive B de R , et n est un lien de la composante de B effacée par $[x]$.

Réciproquement, si n n'a pas de résidus dans R' , alors on a les possibilités suivantes :

- (i) si n est un lien logique, alors (a) ou (b) est satisfaite, ou bien x est un lien coupure logique et la conclusion de n est une prémisse de x .
- (ii) si n est un lien axiome, alors (a) ou (b) est satisfaite, ou bien $[x]$ est une e.e.r. axiome et n est le lien axiome effacé par $[x]$.
- (iii) si n est un lien affaiblissement, alors (a) ou (b) est satisfaite, ou bien la conclusion de n est une prémisse du lien coupure affaiblissement x et la boîte exponentielle effacée par $[x]$ n'a pas de portes auxiliaires.
- (iv) si n est un lien coupure, alors (a) ou (b) est satisfaite, ou bien $n = x$ est un lien coupure logique, ou bien $n = x$ est un lien coupure affaiblissement, ou bien $n = x$ et $[x]$ est une e.e.r. axiome.

2.2.11. REMARQUE. Soient R et R' deux réseaux t.q. $R \xrightarrow{[x]} R'$, et considérons (avec les notations de la remarque précédente) le cas où n est un lien avec (c.à.d. une boîte additive) ou un lien bien sûr (c.à.d. une boîte exponentielle). Soit B la boîte associée à n .

(i) Si B est une boîte additive de R et si il n'y a pas de résidu de B dans R' , alors on a exactement trois possibilités :

- $[x] = [w]$ est une e.e.r. qui efface une boîte exponentielle contenant la boîte additive B
- $[x] = [\&/\oplus]$ est une e.e.r. logique qui efface un sous-réseau de R contenant la boîte additive B
- $[x] = [\&/\oplus]$, où une prémissse du lien coupure $x = \&/\oplus$ est la porte $\&$ de B .

Dans les deux premiers cas nous dirons que B est effacée par $[x]$, tandis que dans le troisième nous dirons que B est ouverte par $[x]$.

(ii) Si $B^!$ est une boîte exponentielle de R qui n'a pas de résidus dans R' , alors on a exactement trois possibilités :

- $[x] = [w]$ est une e.e.r. affaiblissement et la boîte $B_1^!$ de R t.q. la conclusion de la pal de $B_1^!$ est une prémissse de x contient $B^!$ (on pourrait avoir $B_1^! = B^!$)
- $[x] = [\&/\oplus]$ est une e.e.r. logique qui efface un sous-réseau de R contenant la boîte exponentielle $B^!$
- x est un lien coupure $!/?de$ dont une prémissse est la conclusion de la pal de $B^!$.

Dans les deux premiers cas nous dirons que B est effacée par $[x]$, tandis que dans le troisième nous dirons que B est ouverte par $[x]$.

2.3 Préservation de la correction

Il s'agit maintenant de montrer que les définitions statiques du chapitre précédent ayant conduit à une représentation des preuves sous forme de réseaux (le théorème 1.3.10), sont compatibles avec la dynamique interne de ces mêmes réseaux (la normalisation).

Dans cette section, nous allons distinguer les graphes de correction dans le sens de la définition 1.2.6 des graphes de correction dans le sens de la définition 1.3.6 : nous appellerons ces derniers graphes de correction complets.

2.3.1. LEMME. Soit x un lien coupure du réseau R , et supposons que $R \xrightarrow{[x]} R'$. R' est alors un préréseau.

Démonstration : Il est bien clair que R' est une préstructure de preuve. Il s'agit donc de vérifier que R' satisfait les quatre conditions de la définition 1.2.3 de pré-réseau :

- (i) est évidente.
- (ii) est clairement satisfaite, modulo un éventuel renommage des variables (à cause des possibles duplications).
- (iii) est une simple vérification : c'est une conséquence du fait qu'aucun lien ne peut "sortir d'une boîte" lors d'une e.e.r..

Il s'agit donc de montrer que (iv) est satisfaite : il faut vérifier que si la variable X apparaît libre dans la formule étiquette d'une conclusion de la boîte B' de R' , alors X n'est la variable propre d'aucun lien $\forall$ de B' . Supposons par absurde que n' soit un lien $\forall$ de R' contenu dans B' et de variable propre X . Puisque R est un pré-réseau, la/les préstructure(s) de preuve contenue(s) dans l'ancêtre B de B' est/sont un/des pré-réseau(x). L'e.e.r. $[x]$ est donc forcément une e.e.r. $[cc]$ ou bien $[ccad]$.

Si $[x] = [cc]$, soit $B_1^!$ la boîte (exponentielle) avalée par B et B_1'' son résidu dans R' . R satisfaisant (iv), c'est forcément une conclusion de $B_1^!$ qui est étiquetée par une formule contenant X libre. Mais ceci contredit (iii) de la définition 1.2.3 de pré-réseau.

Si $[x] = [ccad]$, soit S le sous-réseau de R avalé par B . R satisfaisant (iv), c'est forcément une conclusion de S qui est étiquetée par une formule contenant X libre. Par (iii) de la définition 1.2.3, n n'est pas un lien de B , donc c'est forcément un lien de S . Mais S est un réseau, donc un pré-réseau, et il ne peut donc pas y avoir dans S un lien $\forall$ dont la variable propre apparaît dans les étiquettes des conclusions de S . $\square$

Dans l'énoncé et dans la démonstration qui suivent, nous ferons référence à la définition préliminaire (1.2.7) de réseau. Le réduit en une étape d'un réseau R est alors parfaitement défini, pourvu que l'étape en question ne soit pas une étape $[ccad]$.

2.3.2. THÉORÈME. Soit x un lien coupure du réseau R , et supposons que $R \xrightarrow{[x]} R'$, avec $[x] \neq [ccad]$. R' est alors un réseau.

Démonstration : Il s'agit de montrer que tous les graphes de correction (dans le sens de la définition 1.2.6) sont acycliques. Nous n'allons pas reproduire ici la preuve de la proposition 6.1 p. 40 de [Danos 90] qui règle la question dans le cas multiplicatif et exponentiel. La présence des additifs (tant que $[x] \neq [ccad]$) est compatible avec l'argument, ainsi que celle des quantificateurs si $[x] \neq [\forall/\exists]$.

Nous allons donc analyser le cas où $[x] = [\forall/\exists]$, en adaptant à PN_2 la preuve de [Girard 91a]. Soient $\forall YA$ et $\exists YA^\perp$ les étiquettes des prémisses de x et soit $A[Y]$ (resp. $A^\perp[B/Y]$) l'étiquette de la prémissse du lien $\forall$ (resp. $\exists$) dont la conclusion est une prémissse de x . Dans ce cas, tous les liens et les arêtes de R (resp. de R') ont un unique résidu (resp. ancêtre) dans R' (resp. dans R), sauf les arêtes étiquetées par $\forall YA$ et $\exists YA^\perp$ et les deux liens dont elles sont conclusions.

Soit G' un graphe de correction de R' . On va montrer que l'acyclicité de tous les graphes de correction de R implique l'acyclicité de G' . On va supposer que x est à profondeur exponentielle et additive zéro.

Si aucune variable libre de B n'est la variable propre d'un lien de R , l'acyclicité de G' découle de celle du graphe de correction G de R , obtenu en "étendant" G' de façon triviale : parmi le k -uple d'arêtes appariées ayant le lien $\forall$ de variable propre Y comme but, on choisit la prémissse du lien.

Autrement, soient $m'_1, \dots, m'_n$ les liens de R' qui sont connectés dans G' respectivement aux liens $\forall$ de R' $\forall'_1, \dots, \forall'_n$, et qui sont tels que dans aucun graphe de correction de R leurs ancêtres $m_1, \dots, m_n$ ne sont connectés aux ancêtres $\forall_1, \dots, \forall_n$ de $\forall'_1, \dots, \forall'_n$. Si aucun tel lien n'existe, on procédera exactement comme ci-dessus (le cas que nous considérons actuellement est bien le seul cas pour lequel on ne puisse pas étendre trivialement G' en un graphe de correction de R).

Remarquons que la variable propre des liens $\forall_i$ est une variable libre de B , et que pour tout $i \in \{1, \dots, n\}$ une prémissse de m_i contient Y libre : nous définissons alors G_Y comme G' sauf pour le lien $\forall$ de variable propre Y (notons-le $\forall_Y$) et pour les liens $\forall_1, \dots, \forall_n$. Pour ces derniers liens on connecte $\forall_i$ au lien $\exists$ de R dont la conclusion est une prémissse de x , et ce pour tout $i \in \{1, \dots, n\}$. Ceci est possible, puisque la variable propre des liens $\forall_i$ est une variable libre de B . Remarquons que G_Y n'est pas un graphe de correction, puisque nous n'avons encore rien dit au sujet du lien $\forall_Y$.

Soit $G^\cap$ le graphe obtenu à partir de G_Y en effaçant les arêtes connectant les liens $\forall_i$ au lien $\exists$ de R dont la conclusion est une prémissse de x . $G^\cap$ peut être vu comme l'intersection de G_Y et de G' . Puisque toute extension de G_Y en un graphe de correction est acyclique, pour tout $i, j \in \{1, \dots, n\}, i \neq j$, $\forall'_i$ et $\forall'_j$ appartiennent à deux composantes connexes différentes de $G^\cap$, et ni l'un ni l'autre n'appartient à la composante connexe de $G^\cap$ qui contient le résidu de x dans R' .

Pour conclure, il suffira de montrer que pour tout $i, j \in \{1, \dots, n\}$ (éventuellement $i = j$), si on appelle C'_j la composante connexe de $\forall'_j$ dans $G^\cap$, alors $m'_i \notin C'_j$. Si $m'_i \in C'_j$, alors C_j (la composante connexe de $\forall_j$ dans $G^\cap$), contient m_i . Mais alors, en choisissant de connecter $\forall_Y$ avec m_i (dont une prémissse,

rappelons-le, contient Y libre), on obtiendrait (à partir de G_Y) un graphe de correction G de R qui est cyclique. $\square$

2.3.3. LEMME. Si (R, j_R) est un réseau (avec sauts), x est un lien coupure de R , $R \xrightarrow{[x]} R'$, et $j_{R'}$ est une des fonctions associées à j_R par la définition 2.2.8, alors $(R', j_{R'})$ est un préréseau avec sauts.

Démonstration : Nous savons que R' est un préréseau (lemme 2.3.1). Il s'agit alors de faire preuve d'un peu de patience, et de prouver que $(R', j_{R'})$ satisfait les conditions (1), (2) et (3) de la définition 1.3.3.

(1) : Soit n' un lien affaiblissement de R' , soit $y' \in j_{R'}(n')$, B' une boîte additive de R' contenant y' et ne contenant pas n' . Appelons n (resp. y , B) l'ancêtre dans R de n' (resp. y' , B'). B ne contient pas n . Si B ne contient pas y , alors (le seul moyen pour entrer dans une boîte additive étant celui d'appliquer une $[ccad]$ attractive pour cette boîte) $[x] = [ccad]$ B -attractive et y est un lien du sous-réseau S avalé par B . Dans ce cas, il existe aussi un résidu y'' de y dans B' , et B' sépare y' et y'' . On peut donc supposer que B contient y et ne contient pas n .

(1.1) $y \in j_R(n)$: alors (par (1) de la définition 1.3.3), B sépare y et un certain $z \in j_R(n)$. Si B a deux résidus dans R' , alors appelons z'_1 et z'_2 les deux résidus de z dans R' . B' sépare alors y' et z'_i , avec $z'_i \in j_{R'}(n')$ pour un certain $i \in \{1, 2\}$.

Si B a un unique résidu dans R' et si z en a (au moins) un, alors B' sépare y' et tout résidu de z , un (au moins) de ces derniers étant un élément de $j_{R'}(n')$.

Si B a un unique résidu dans R' et z n'en a aucun, alors x est contenu dans la composante $\gamma_i(B)$ de B qui ne contient pas y , et (par la remarque 2.2.10) $[x] = [w]$, $[x] = [\&/\oplus]$, ou $[x] = [ax]$. Dans les cas $[x] = [w]$ et $[x] = [ax]$, la définition 2.2.8 garantit l'existence d'au moins un lien $t' \in j_{R'}(n')$ t.q. B' sépare y' et t' . Si $[x] = [\&/\oplus]$, appelons B_1 la boîte additive de R ouverte par $[x]$: z est un lien de la composante de B_1 effacée par $[x]$. Mais puisque (par (1) de la définition 1.3.3) B_1 sépare z et un certain $t \in j_R(n)$, t a un résidu t' dans R' et B' sépare y' et $t' \in j_{R'}(n')$.

(1.2) $y \notin j_R(n)$: alors par la remarque 2.2.9, $[x] = [w]$ ou bien $[x] = [ax]$. Supposons pour fixer les idées que $y \in \gamma_1(B)$.

(1.2.1) Il existe $z \in \gamma_2(B)$ t.q. $z \in j_R(n)$: Si z a (au moins) un résidu z' dans R' , alors par la remarque 2.2.9 $z' \in j_{R'}(n')$. B' sépare alors y' et z' . Si z n'a pas de résidus dans R' , alors x est un lien de $\gamma_2(B)$, et

(comme ci-dessus) la définition 2.2.8 garantit l'existence d'au moins un lien $t' \in j_{R'}(n')$ t.q. B' sépare y' et t' .

(1.2.2) Aucun lien de $\gamma_2(B)$ n'est élément de $j_R(n)$:

(1.2.2.1) $[x] = [ax]$: B est alors une des boîtes additives satisfaisant (i) du lemme 2.1.5, et il existe donc z t.q. B sépare y et z et (par la définition 2.2.8) $y', z' \in j_{R'}(n')$. B' sépare dans ce cas y' et z' .

(1.2.2.2) $[x] = [w]$: on a alors $y \in j_R(m)$, où m est le lien affaiblissement dont la conclusion est prémisses de x . Si $m = n$, alors (par (1) de la définition 1.3.3) il existe $z \in j_R(n)$ t.q. B sépare y et z , c.à.d. que $z \in \gamma_2(B)$, ce qu'on a exclu.

Donc $m \neq n$ et il existe un lien t contenu dans la plus grande boîte exponentielle $B^!$ effacée par x t.q. $t \in j_R(n)$. Si m est contenu dans B , alors (puisque $y \in j_R(m)$) m est forcément un lien de $\gamma_1(B)$, et $B^!$ est contenue dans $\gamma_1(B)$. Mais on sait que n n'est pas contenu dans B et que $t \in j_R(n)$ est contenu dans $\gamma_1(B)$: toujours par (1) de la définition 1.3.3, B sépare donc t et un certain $z \in j_R(n)$, avec $z \in \gamma_2(B)$, ce qu'on a exclu. Donc m n'est pas contenu dans B . Mais alors $y \in j_R(m)$, m n'est pas contenu dans B et y est contenu dans B : B sépare y et un certain $z \in j_R(m)$. Alors B' sépare y' et le résidu de z , qui est un élément de $j_{R'}(n')$ (par la définition 2.2.8).

(2) : Soient $y', z' \in j_{R'}(n')$, et soient y et z les ancêtres de y' et z' dans R . Si $y = z$, alors forcément $[x] = [ccad]$ ou bien $[x] = [co]$. Le cas $[x] = [co]$ est en fait exclus par la définition 2.2.8 : on aurait dans ce cas $y' \notin j_{R'}(n')$ ou bien $z' \notin j_{R'}(n')$. Donc $[x]$ est une e.e.r. $[ccad]$ B -attractive pour une certaine boîte B de R , dont le résidu dans R' sépare y' et z' . On peut donc supposer $y' \neq z'$.

(2.1) $y \in j_R(n)$ et $z \in j_R(n)$: Soit B_{yz} la boîte de R qui sépare y et z . Puisque y et z ont tous les deux (au moins) un résidu dans R' , B_{yz} a aussi au moins un résidu dans R' . Si B_{yz} a un unique résidu, alors celui-ci sépare y' et z' . Autrement, $[x] = [co]$ ou bien $[x] = [ccad]$. Si $[x] = [co]$, la définition 2.2.8 implique que un des deux résidus de B_{yz} sépare y' (*l'unique* résidu de y qui est un saut de n') et z' (*l'unique* résidu de z qui est un saut de n'). Si $[x] = [ccad]$ B -attractive, et si aucun des deux résidus de B_{yz} ne sépare y' et z' , alors c'est le résidu de B qui sépare y' et z' .

(2.2) $y \notin j_R(n)$: alors $j_{R'}(n')$ ne contient pas seulement des résidus d'éléments de $j_R(n)$, et la remarque 2.2.9 nous dit qu'alors $[x] = [ax]$ ou bien $[x] = [w]$.

(2.2.1) $[x] = [ax]$: soit p le lien axiome effacé par $[x]$. $j_{R'}(n')$ ne contenant pas seulement des résidus d'éléments de $j_R(n)$, on a forcément $p \in j_R(n)$. Si $z \in j_R(n)$, alors il existe B séparant p et z dans R . L'unique résidu de B dans R' sépare y' et z' .

Si $z \notin j_R(n)$, alors la définition 2.2.8 (en fait c'est plutôt le lemme 2.1.5 qu'on utilise ici) garantit l'existence d'une boîte B de R qui sépare y et z . L'unique résidu de B dans R' sépare donc y' et z' .

(2.2.2) $[x] = [w]$: soit p le lien affaiblissement dont la conclusion est prémisses de x , et soit B^1 la plus grande boîte exponentielle de R effacée par $[x]$. Ni y ni z n'est contenu dans B^1 . Puisque par ailleurs $y \notin j_R(n)$, on a forcément $y \in j_R(p)$.

Si $z \in j_R(p)$, alors (par (2) de la définition 1.3.3) il existe une boîte additive B de R qui sépare y et z : l'unique résidu de B dans R' sépare y' et z' .

Si $z \notin j_R(p)$, alors $n \neq p$, $z \in j_R(n)$ et il existe un lien t contenu dans B^1 t.q. $t \in j_R(n)$. Soit alors B la boîte de R qui sépare t et z . Comme z n'est pas contenu dans B^1 , B sépare aussi p et z , et donc y et z : l'unique résidu de B dans R' sépare y' et z' .

(3) : Supposons par absurde qu'il existe une boîte B' de R' contenant n' , et un lien $y' \in j_{R'}(n')$ qui n'est pas contenu dans B' . Appelons n (resp. y , B) l'ancêtre de n' (resp. y' , B') dans R . Le lien y n'est pas contenu dans B . Distinguons deux cas :

(3.1) n est un lien de B : puisque (R, j_R) satisfait la condition (3) de la définition 1.3.3, on a $y \notin j_R(n)$. Par la remarque 2.2.9, on a alors $[x] = [ax]$ ou bien $[x] = [w]$: dans les deux cas, la définition 2.2.8 impose (pour que $y' \in j_{R'}(n')$) que y soit à profondeur additive supérieure ou égale à celle de n , et ce n'est pas le cas ici.

(3.2) n n'est pas un lien de B : alors forcément $[x] = [ccad]$ B -attractive. Appelons (S, j_S) le sous-réseau dupliqué par $[x]$. n est un lien de S et par définition de sous-réseau $j_R(n)$ ne contient que des liens de S . Par la remarque 2.2.9, $j_{R'}(n')$ ne contient dans ce cas que des résidus de liens de $j_R(n)$. Donc tous les liens de $j_{R'}(n')$ doivent être contenus dans B' . Or y' n'est pas contenu dans B' .

□

Le lemme suivant est un résultat très simple de théorie des graphes que nous utiliserons dans la preuve du théorème qui suit. Nous ne priverons pas le lecteur du plaisir de retrouver tout seul la démonstration, par exemple par induction sur le nombre de sommets.

2.3.4. LEMME. Soit G un graphe non orienté (et non apparié). Si G est acyclique, alors on a l'égalité suivante :

nombre de sommets-nombres d'arêtes=nombre de composantes connexes.

2.3.5. THÉORÈME. Si (R, j_R) est un réseau, x est un lien coupure de R , $R \xrightarrow{[x]} R'$, et $j_{R'}$ est une des fonctions associées à j_R par la définition 2.2.8, alors $(R', j_{R'})$ est un réseau.

Démonstration : Nous savons par le lemme 2.3.3 que $(R', j_{R'})$ est un pré-réseau avec sauts.

On va d'abord montrer que tous les graphes de correction complets de $(R', j_{R'})$ sont acycliques (y compris bien sûr les graphes de correction complets des contenus des boîtes de $(R', j_{R'})$). Par le théorème 2.3.2, nous savons que c'est le cas des graphes de correction (non complets) lorsque $[x] \neq [ccad]$, et le lecteur n'aura aucun mal à se convaincre que ceci reste vrai même dans le cas $[x] = [ccad]$. Il s'agit donc de montrer que le fait de rajouter, dans un graphe de correction, les arêtes connectant chaque affaiblissement avec le saut qui lui est associé dans le graphe de correction complet, ne va pas créer de cycles.

Si pour tout lien affaiblissement n' de R' l'ensemble $j_{R'}(n')$ ne contient que des résidus de liens de $j_R(n)$ (où n est l'ancêtre de n' dans R), alors il s'agit d'adapter l'argument de la preuve du théorème 2.3.2 aux graphes de correction complets : on montre qu'à partir d'un graphe de correction complet et cyclique de R' on peut construire un graphe de correction complet et cyclique de R . La vérification se fait sans surprises.

Autrement, par la remarque 2.2.9, $[x] = [ax]$ ou bien $[x] = [w]$. Nous allons traiter le cas $[x] = [ax]$ (l'autre étant similaire). Soit p le lien axiome de R effacé par $[x]$ et soient $n_1, \dots, n_k$ les $k \geq 0$ liens affaiblissement de R t.q. $\forall i \in \{1, \dots, k\}$ on a $j_R(n_i) = \{p\}$. Soit T le plus grand sous-réseau de R contenant $n_1, \dots, n_k$ et p . On a : $T \xrightarrow{[x]} T'$, où T' est un sous-réseau de R' . Si n' est un lien affaiblissement de T' t.q. $j_{R'}(n')$ ne contient pas seulement des résidus de liens de $j_R(n)$ (où n est l'ancêtre de n' dans R), alors n' est l'un des résidus $n'_1, \dots, n'_k$ de $n_1, \dots, n_k$.

Soit G' un graphe de correction complet de T' . On va montrer que G' est acyclique.

Pour tout $i \in \{1, \dots, k\}$, appelons l'_i le lien hypothèse de T'_{ap} qui subsitue la plus grande boîte de T' contenant tous les liens de $j_{R'}(n'_i)$ si elle existe, et le seul lien axiome de $j_{R'}(n'_i)$ autrement. Soient $l_1, \dots, l_k$ les ancêtres dans T_{ap} des liens $l'_1, \dots, l'_k$ (même dans le cas où l'_i est un lien hypothèse, le lien l_i est clairement défini). Soit $G'^{\cap}$ le sous-graphe de G' obtenu en effaçant toutes les arêtes connectant n'_i avec l'_i ($i \in \{1, \dots, k\}$). Il existe un graphe de correction complet G de T t.q. en effaçant les arêtes connectant $n_1, \dots, n_k$ à p , et en substituant le lien p et ses conclusions par une arête, on obtient $G'^{\cap}$.

Appelons $G^\cap$ le sous-graphe de G obtenu en effaçant les arêtes connectant $n_1, \dots, n_k$ à p . Puisque les graphes de correction de T sont acycliques et connexes, $G^\cap$ l'est aussi, et les $k + 1$ liens $l_1, \dots, l_k, p$ sont dans la même composante connexe de $G^\cap$. On a alors :

- (1) $\forall i, j \in \{1, \dots, k\}, i \neq j$, n'_i et n'_j sont dans deux composantes connexes différentes de $G'^\cap$: dans le cas contraire n_i et n_j seraient dans la même composante connexe de $G^\cap$, et G serait cyclique
- (2) $\forall i, j \in \{1, \dots, k\}$, n'_i et l'_j sont dans deux composantes connexes différentes de $G'^\cap$: dans le cas contraire n_i et l_j (donc n_i et p) seraient dans la même composante connexe de $G^\cap$, et G serait cyclique.

Il est bien clair que (1) et (2) suffisent pour conclure que G' est acyclique.

Pour montrer que tous les graphes de correction complets de $(R', j_{R'})$ sont connexes, on va utiliser le lemme 2.3.4. On rajoute aux conclusions des graphes appariés des noeuds terminaux (de sorte que les graphes de correction complets soient des graphes dans le sens usuel de la théorie des graphes). Il s'agit ensuite de vérifier (et nous laisserons ce soin au lecteur) que toute e.e.r. maintient constante, à toute profondeur et dans tout graphe de correction complet, l'entier suivant :

nombre de sommets-nombres d'arêtes.

Celui-ci étant égal à un dans tous les graphes de correction complets de (R, j_R) (par connexité et par le lemme 2.3.4), il va demeurer égal à un dans tous les graphes de correction complets de $(R', j_{R'})$: en appliquant de nouveau le lemme 2.3.4 on aura la connexité de ces derniers. $\square$

Nous allons conclure le chapitre en introduisant les notions d'ancêtre et de résidu pour les conclusions et les sous-réseaux d'un réseau.

2.3.6. DÉFINITION. Soit R' un réduit en une étape du réseau R . Il existe une bijection évidente $i_{R,R'}$ entre les conclusions de R et celles de R' . Si $i_{R,R'}(a) = a'$, on dira que l'arête a (conclusion de R) est **l'ancêtre** de l'arête a' (conclusion de R') et que a' est **le résidu** de a dans R' .

La notion d'ancêtre et de résidu est légèrement plus délicate pour les sous-réseaux que pour les liens. Soit $R \xrightarrow{[x]} R'$. Puisque lors de l'élimination des coupures un sous-réseau peut être effacé (par une e.e.r. $[\&/\oplus]$ ou $[w]$) ou modifié (par une e.e.r. $[x]$ t.q. x est un lien coupure qui appartient au sous-réseau), un sous-réseau α de R' provient d'*au plus* une occurrence du "même" (à des substitutions près) sous-réseau $\overleftarrow{\alpha}$ de R . On aurait envie de dire qu'un

tel sous-réseau de R (si il existe) est l'ancêtre de α dans R . Mais une telle définition serait ambiguë : soit $B^!$ une boîte exponentielle de R' , $\overleftarrow{B}^!$ son ancêtre dans R , supposons que x soit un lien coupure contraction et que la conclusion de la pal de $\overleftarrow{B}^!$ soit une prémissse de x . Supposons que le sous-graphe α de R' qui consiste en la boîte $B^!$ et le lien affaiblissement n (qui n'appartient pas à $B^!$) avec sa conclusion soit un sous-réseau de R' (on dit tout simplement que tous les liens de $j_{R'}(n)$ sont des liens de $B^!$). Allons-nous considérer le sous-réseau de R qui consiste en $\overleftarrow{B}^!$ et l'ancêtre $\overleftarrow{n}$ de n (avec sa conclusion) comme l'ancêtre de α ? (Remarquons que $\overleftarrow{B}^!$ a deux résidus tandis que $\overleftarrow{n}$ a un seul résidu dans R'). La réponse que fournit la définition suivante est négative.

2.3.7. DÉFINITION. Soient R et R' deux réseaux t.q. $R \xrightarrow{[x]} R'$, et soit α un sous-réseau de R' . α provient d'au plus une occurrence du "même" (à des substitutions près si $[x] = [\forall/\exists]$) sous-réseau $\overleftarrow{\alpha}$ de R . Si un tel sous-réseau $\overleftarrow{\alpha}$ existe, nous dirons que c'est **l'ancêtre du sous-réseau** α lorsque une (et une seule) des deux conditions suivantes est satisfaite :

- tout lien principal de α a un unique résidu dans R'
- tout lien principal de α a deux résidus dans R' .

Réciproquement, on définit l'ensemble des **résidus du sous-réseau** β de R , comme l'ensemble de toutes les occurrences β' de sous-réseaux de R' telles que β est l'ancêtre de β' .

C'est dans la remarque suivante que réside en fait la vraie motivation pour la définition précédente.

2.3.8. REMARQUE. Soient T et T' deux réseaux t.q. $T \xrightarrow{[x]} T'$, où x est un lien coupure contraction. Soit n le lien bien sûr de T dont la conclusion est une prémissse de x et soit $B^!$ la boîte exponentielle de T associée à n . Si R est un sous-réseau de T qui a un résidu dans T' et si R contient $B^!$, alors $R = B^!$.

2.3.9. REMARQUE. (i) Si σ est une suite d'e.e.r. issue de (R, j_R) , alors on peut définir la notion de σ -réduit (ou simplement réduit) de R . Si $(R', j_{R'})$ est un σ -réduit de R , alors on écrit $(R, j_R) \xrightarrow{\sigma} (R', j_{R'})$. En fait, on écrira souvent plutôt $R \xrightarrow{\sigma} R'$, ou simplement $R \rightarrow R'$.

(ii) Les notions d'ancêtre et de résidu d'un lien logique, axiome, coupure ou affaiblissement, ainsi que celles d'ancêtre et de résidu d'une boîte, d'un sous-réseau ou d'une conclusion se généralisent immédiatement à une suite d'e.e.r.

(iii) Nous insistons ici sur le fait qu'un σ -réduit $(R', j_{R'})$ n'est pas univoquement déterminé par (R, j_R) et σ ; la notation $(R, j_R) \xrightarrow{\sigma} (R', j_{R'})$ indique simplement qu'il *existe* une façon d'obtenir $(R', j_{R'})$ à partir de (R, j_R) en appliquant aux e.e.r. de σ la définition 2.2.8. En particulier, lorsque nous ferons usage de la notion de substitution d'un réseau à un sous-réseau dans un réseau donné (définie dans 2.1.6), il nous arrivera d'écrire que (sous certaines hypothèses bien sûr) si $T \xrightarrow{\sigma} T'$, alors $T[\beta/\alpha] \xrightarrow{\sigma} T'[\beta/\alpha'_1, \dots, \beta/\alpha'_n]$, où α est un sous-réseau de T , β un réseau avec les mêmes conclusions que α , et $\alpha'_1, \dots, \alpha'_n$ sont les résidus de α dans T' (cf. définition 2.3.7). On veut dire par là qu'il existe une fonction de saut $j_{T[\beta/\alpha]}$, une fonction de saut $j_{T'[\beta/\alpha'_1, \dots, \beta/\alpha'_n]}$, et une façon d'obtenir $(T'[\beta/\alpha'_1, \dots, \beta/\alpha'_n], j_{T'[\beta/\alpha'_1, \dots, \beta/\alpha'_n]})$ à partir de $(T[\beta/\alpha], j_{T[\beta/\alpha]})$ en appliquant aux e.e.r. de σ la définition 2.2.8.

Chapitre 3

La propriété de séparation

Ce court chapitre est consacré à une propriété de la normalisation spécifique à la présence du connecteur $\&$, dont il sera de nouveau question dans le chapitre 5. Cette propriété de séparation (le théorème 3.2.2) est essentielle dans la preuve du théorème de standardisation additive de l'Annexe A.

3.1 Le lemme de substitution

Le lemme suivant va se révéler très utile dans la suite de la thèse.

3.1.1. LEMME. (Lemme de substitution) Soient T, T' deux réseaux, t.q. $T \xrightarrow{\sigma} T'$. Soit α un sous-réseau de T , et β un réseau de mêmes conclusions que α .

Si, pour tout $[x] \in \sigma$ t.q. $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T_2 \xrightarrow{\sigma_2} T'$, et pour tout résidu γ_1 de α dans T_1 , les conditions suivantes sont satisfaites :

- i) si $[x] = [ax]$ alors x n'est pas contenu dans γ_1 et le lien ax effacé par $[x]$ n'est pas contenu dans γ_1
- ii) si $[x] \neq [ax], [ccad]$ alors aucune des (deux) prémisses de x n'est une arête de γ_1 (en particulier, toute conclusion de γ_1 n'est pas prémisses de x)
- iii) si $[x] = [ccad]$ alors x n'est pas un lien de γ_1 , $[x]$ n'est attractive pour aucune boîte de γ_1 et si $[x]$ duplique un sous-réseau de γ_1 alors $[x]$ duplique γ_1 tout entier

alors, si $\alpha'_1, \dots, \alpha'_n$ sont les résidus de α dans T' , on a $T[\beta/\alpha] \xrightarrow{\sigma} T'[\beta/\alpha'_1, \dots, \beta/\alpha'_n]$ (à des substitutions près). Si $[\forall/\exists] \in \sigma$, alors on aura, plus précisément, $T[\beta/\alpha] \xrightarrow{\sigma} T'[\beta_1/\alpha'_1, \dots, \beta_n/\alpha'_n]$, où $\beta_i = \beta$ dans lequel on a effectué les substitutions requises par la ou les e.e.r. $[\forall/\exists]$.

Démonstration : Intuitivement, les conditions (i), (ii), (iii) garantissent que $[x]$ ne modifie rien dans les résidus de α : la seule chose que $[x]$ puisse faire c'est dupliquer ou effacer des résidus.

Plus formellement, nous procédons par induction sur σ . Si σ est vide, alors $T = T'$ et le résultat est évident.

Autrement $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$. Par hypothèse d'induction, si $\alpha_1^1 \dots \alpha_1^m$ sont les résidus de α dans T_1 , on a $T[\beta/\alpha] \xrightarrow{\sigma_1} T_1[\beta/\alpha_1^1, \dots, \beta/\alpha_1^m]$ (à des substitutions près). Pour conclure, il suffit de vérifier que pour tout résidu α_1^i de α dans T_1 et pour tous les résidus α' de α_1^i dans T' , on a $T_1[\beta/\alpha_1^i] \xrightarrow{[x]} T'[\beta/\alpha']$ (à des substitutions près si $[x] = [\forall/\exists]$), où $[x]$ est une quelconque e.e.r. satisfaisant les conditions i), ii) et iii). Dans le cas où $[x]$ est une e.e.r. $[ccad]$ qui duplique le sous-réseau δ de T_1 , il est important de remarquer que (grâce à (iii) de la remarque 2.1.7) $\delta[\beta/\alpha_1^1, \dots, \beta/\alpha_1^m]$ est un sous-réseau de $T_1[\beta/\alpha_1^1, \dots, \beta/\alpha_1^m]$. $\square$

3.2 Normalisation et séparation

La remarque suivante est aussi simple que fondamentale.

3.2.1. REMARQUE. Soit R un réseau et soit R' un réduit de R . Si l est un lien logique, un lien axiome ou un lien coupure *logique*, de profondeur exponentielle zéro dans R , alors tout résidu de l dans R' est à profondeur exponentielle zéro dans R' . Ceci reste vrai pour un sous-réseau α de R différent d'une boîte exponentielle, et pour une boîte additive B de R : si α et B sont à profondeur exponentielle zéro dans R , alors tout résidu de α dans R' ainsi que tout résidu de B dans R' est à profondeur exponentielle zéro dans R' .

Ceci parce que la seule possibilité pour un lien (différent d'un lien coupure commutative exponentielle) d'entrer dans une boîte exponentielle est *d'être déjà* dans une autre boîte exponentielle, dont la conclusion de la pal est prémisses d'un lien coupure commutative exponentielle.

Nous allons maintenant prouver une propriété importante de la normalisation en présence d'additifs. Si on suit l'intuition (corroborée par la sémantique dénotationnelle cohérente) qu'une boîte additive n'est autre qu'une représentation syntaxique de la superposition de deux processus calculatoires, le théorème de séparation peut être informellement énoncé de la façon suivante : «la normalisation préserve la superposition à *profondeur exponentielle zéro*».

3.2.2. THÉORÈME. (Propriété de séparation). Soient T, T' deux réseaux, t.q. $T \xrightarrow{\sigma} T'$. Soit $\mathbb{B}$ une boîte additive de profondeur exponentielle zéro dans T , et $\mathbb{B}'$ un résidu de $\mathbb{B}$ dans T' .

Pour tout résidu $\mathbb{B}''$ de $\mathbb{B}$ dans T' t.q. $\mathbb{B}'' \neq \mathbb{B}'$, il existe une boîte additive B de T' qui sépare $\mathbb{B}'$ et $\mathbb{B}''$.

Démonstration : Par induction sur σ . Si $T' = T$, alors le résultat est évident.

Autrement $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$.

Si $\mathbb{B}'$ est l'unique résidu de $\mathbb{B}$ dans T' , alors le lemme est trivialement vrai. Autrement soit $\mathbb{B}'' \neq \mathbb{B}'$ un résidu de $\mathbb{B}$ dans T' . Soit $\mathbb{B}'_1$ (resp. $\mathbb{B}''_1$) l'ancêtre de $\mathbb{B}'$ (resp. $\mathbb{B}''$) dans T_1 .

Si $\mathbb{B}'_1 = \mathbb{B}''_1 = \mathbb{B}_1$, alors (puisque $\mathbb{B}'' \neq \mathbb{B}'$) $\mathbb{B}_1$ est dupliquée lors de l'e.e.r. $[x]$. Mais comme $\mathbb{B}_1$ est à profondeur exponentielle zéro dans T_1 (à cause de la remarque 3.2.1), la seule possibilité est que $[x] = [ccad]$. Soit alors B_1 la boîte additive de T_1 t.q. $[x]$ est B_1 -attractive, et soit S le sous-réseau de T_1 avalé par B_1 lors de l'e.e.r. $[x]$. $\mathbb{B}_1$ est (évidemment) un sous-réseau de S . La boîte B de T' cherchée est alors l'unique résidu de B_1 dans T' .

Supposons maintenant que $\mathbb{B}'_1 \neq \mathbb{B}''_1$. Par hypothèse d'induction, il existe une boîte additive B_1 de T_1 qui sépare $\mathbb{B}'_1$ et $\mathbb{B}''_1$.

Remarquons que la boîte B_1 de T_1 a un ou deux résidu(s) dans T' : si B_1 n'avait pas de résidu dans T' , alors (voir aussi (ii) de la remarque 2.2.11) une au moins des deux composantes de B_1 disparaîtrait, et alors une boîte parmi $\mathbb{B}'_1$ et $\mathbb{B}''_1$ n'aurait pas de résidu dans T' , ce qui est impossible.

Enfin, considérons les deux possibilités :

- (i) B_1 a un résidu dans T' : dans ce cas la boîte B de T' cherchée est l'unique résidu de B_1 dans T' .
- (ii) B_1 a deux résidus dans T' : appelons-les B'_{11} et B'_{12} . Toujours parce que B_1 est à profondeur exponentielle zéro, $[x] = [ccad]$ et B_1 est contenue dans le sous-réseau S de T_1 dupliqué par l'e.e.r. $[x]$. Soit B_2 la boîte additive de T_1 t.q. $[x]$ est B_2 -attractive, et soit B'_2 l'unique résidu de B_2 dans T' .

En tant que sous-réseaux de B_1 (donc de S), $\mathbb{B}'_1$ et $\mathbb{B}''_1$ ont chacun deux résidus dans T' . Plus précisément, un résidu de $\mathbb{B}'_1$ (resp. de $\mathbb{B}''_1$) est un sous-réseau de B'_{11} et l'autre est un sous-réseau de B'_{12} . Appelons $\mathbb{B}'_{11}$ (resp. $\mathbb{B}''_{11}$) le résidu de $\mathbb{B}'_1$ (resp. de $\mathbb{B}''_1$) contenu dans B'_{11} et $\mathbb{B}'_{12}$ (resp. $\mathbb{B}''_{12}$) celui qui est contenu dans B'_{12} . On a $\mathbb{B}' = \mathbb{B}'_{1i}$ et $\mathbb{B}'' = \mathbb{B}''_{1j}$, avec $i, j \in \{1, 2\}$.

Si $i = j = 1$ (resp. $i = j = 2$), alors la boîte B'_{11} (resp. B'_{12}) est la boîte additive de T' cherchée. Si $i \neq j$, alors la boîte B'_2 est la boîte cherchée.

□

3.2.3. REMARQUE. (i) Comme toute boîte qui sépare des liens ou des boîtes, la boîte dont le théorème précédent affirme l'existence est unique.

(ii) Une boîte $\mathbb{B}$ qui ne satisfait pas l'hypothèse “ $\mathbb{B}$ à profondeur exponentielle zéro dans T ”, n'est pas assurée de satisfaire la propriété de séparation.

Supposons par exemple que $T \xrightarrow{[co]} T'$ et que $\mathbb{B}$ soit contenue dans la boîte exponentielle dupliquée par l'e.e.r. $[co]$. Les deux résidus $\mathbb{B}'$ et $\mathbb{B}''$ de $\mathbb{B}$ dans T' ne sont séparés par aucune boîte additive de T' . Ceci contredit la conclusion du théorème, et aussi (ce qui est pire) sa signification intuitive : $\mathbb{B}'$ et $\mathbb{B}''$ ne sont pas superposées dans T' , donc l'étape de normalization $[co]$ n'a pas préservé la superposition.

Mentionnons au passage le fait que ce phénomène semble être lié à l'isomorphisme d'espaces cohérents $!(A \& B)$ et $!A \otimes !B$: «un lien $\&$ à profondeur exponentielle supérieure à zéro n'est pas un vrai lien $\&$ ».

(iii) La propriété de séparation nous dit aussi que, dans tout réduit d'un réseau T , deux résidus différents d'une boîte additive de profondeur exponentielle zéro dans T *ne peuvent pas interagir* (il n'appartiendront jamais, par exemple, à deux sous-réseaux connectés par un lien $\otimes$). Remarquons que ce n'est pas du tout le cas des résidus d'une boîte exponentielle. Intuitivement, ceci veut dire que le phénomène de duplication présent lors d'une e.e.r. commutative additive est de nature complètement différente de celui qui accompagne une e.e.r. contraction.

(iv) Une conséquence de la propriété de séparation est qu'un résidu $\mathbb{B}'$ d'une boîte additive $\mathbb{B}$ de profondeur exponentielle zéro dans T ne peut pas être contenu dans un résidu $\mathbb{B}''$ de $\mathbb{B}$ différent de $\mathbb{B}'$. Ceci est faux en général (si $\mathbb{B}$ n'est pas à profondeur exponentielle zéro) : comme les boîtes exponentielles, les boîtes additives peuvent (en un certain sens) “se rentrer dans elles-mêmes”, ce qui a pour conséquence (notamment) qu'il est impossible de borner à priori la profondeur additive (de même que la profondeur exponentielle) des réduits d'un réseau. Il n'est pas inintéressant de remarquer, que ceci reste vrai dans les systèmes à “basse” complexité tels que *ELL* et *LLL* : si il est vrai que la profondeur exponentielle est constante, les contenus des boîtes exponentielles, eux, peuvent toutefois interagir, de sorte que le phénomène d'une boîte *additive* “s'avalant elle-même” n'aura pas disparu. Dans [Girard 95b], ceci est soigneusement évité : on n'effectue pas les e.e.r. $[ccad]$, ce qui, comme nous allons le voir dans le prochain chapitre, est le plus souvent très raisonnable...

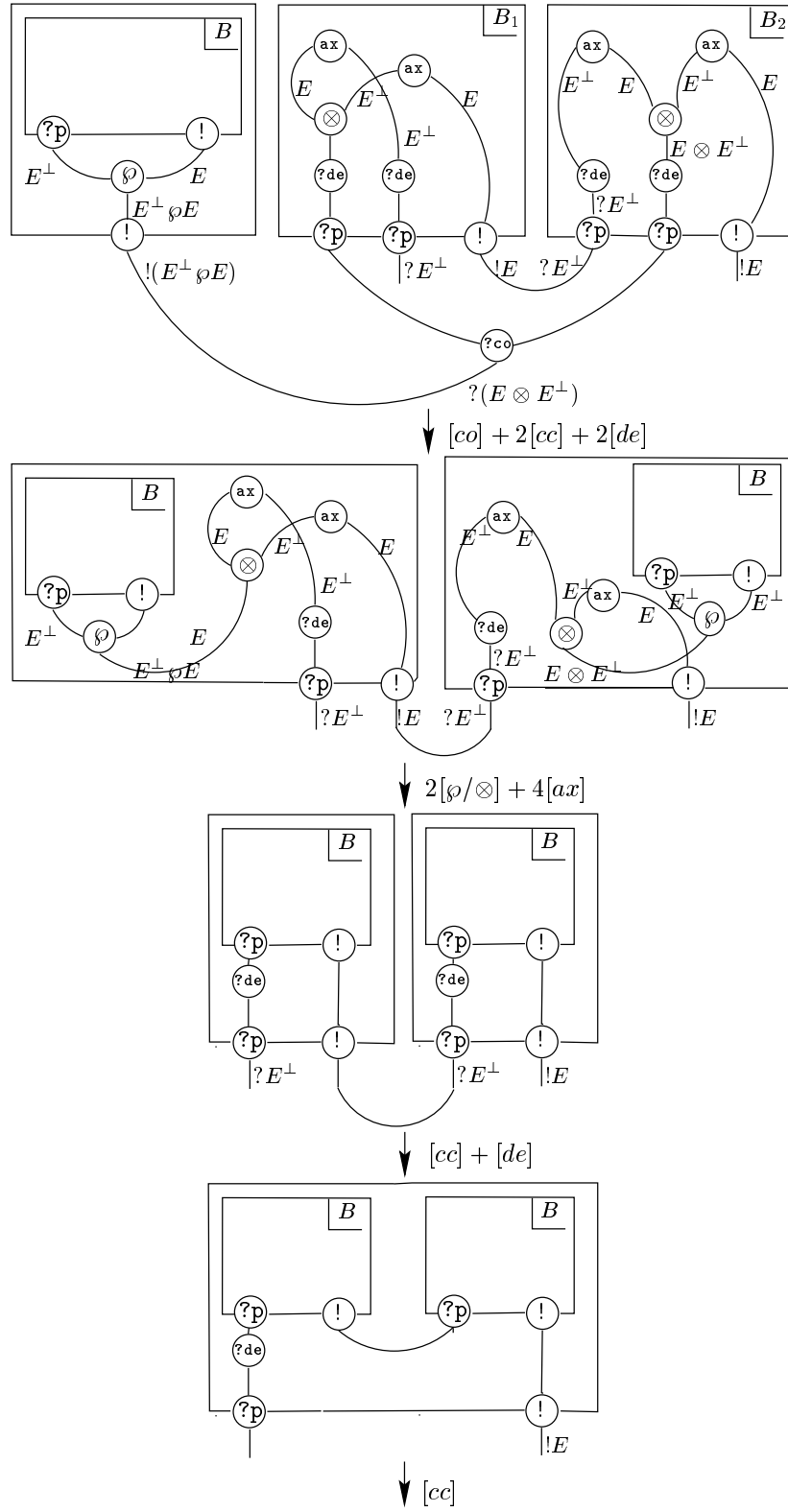
La figure 5 fournit un exemple de boîte exponentielle de PN_2 “s'avalant

elle-même". L'exemple est un peu plus compliqué que nécessaire, pour que le lecteur connaissant le système ELL puisse reconnaître que, malgré le fait que la profondeur exponentielle est constante dans ELL , une boîte *additive* contenue dans la boîte exponentielle B de la figure 5 peut avaler un de ses résidus.

Remarquons aussi que le réseau de départ de la figure 5 correspond au λ -terme $(\lambda f(f))(f)x\lambda yt$: le phénomène en question apparaît alors très clairement.

Le lemme suivant peut s'obtenir par une démonstration presque identique à celle du théorème de séparation, *mutatis mutandis*.

3.2.4. LEMME. Soit T un réseau t.q. $T \xrightarrow{\sigma} T'$, c un lien coupure de T , et supposons que si R est un réseau apparaissant dans σ , alors tout résidu de c dans R est à profondeur exponentielle zéro dans R (ce qui signifie, en particulier, que c est à profondeur exponentielle zéro dans T). Soient c' , c'' deux résidus de c dans T' t.q. $c' \neq c''$. Il existe alors une boîte additive B' qui sépare c'' et c' .

Figure 5. Un résidu de B avale un autre résidu de B

Chapitre 4

Confluence faible des réseaux de PN_2

Le système PN_2 muni de sa normalisation ne jouit pas de la propriété de confluence (ou Church-Rosser). Pour s'en convaincre sans difficultés, le lecteur peut éliminer le lien coupure entre la boîte additive (contenant deux liens axiome) de conclusions $\vdash A^\perp \& A^\perp, A$ et la boîte additive (contenant deux liens axiome) de conclusions $\vdash A \& A, A^\perp$. Le problème est ici que les prémisses du lien coupure sont toutes les deux conclusions de portes coad : il faut choisir (arbitrairement) laquelle des deux boîtes avale (et duplique) l'autre (voir figure 6). Remarquons au passage que ce même conflit est déjà présent avec la définition originelle de $[ccad]$ de [Girard 87].

Cependant, la propriété de séparation semble indiquer que cette non confluence n'est pas, pour ainsi dire, incontrôlable. A cette affirmation informelle, le présent chapitre donne le contenu suivant : si parmi les conclusions du réseau R il n'y a aucune formule contenant une occurrence du connecteur $\&$ ou du quantificateur $\exists$, alors la normalisation définie dans le chapitre 2 associe à R un unique réseau sans coupures (théorème 4.3.5). Signalons que ce cas est très important pour “programmer” avec des réseaux (suivant l'approche décrite dans [KriPar 90] et dite de “programmation par preuves”). Les types des données sont en effet des formules sans $\&$ ni $\exists$, même si ces connecteurs (et les e.e.r. qui leurs sont associées) apparaissent dans le déroulement du calcul (voir aussi la remarque 4.2.3, et [Girard 95a], [DJ 99]).

Dans *ce chapitre*, nous allons considérer les ensembles de réseaux suivants :
- PN_2 : l'ensemble des réseaux avec sauts définis dans 1.3.7, dont les éléments seront appelés réseaux avec sauts

- PN_2^+ : l'ensemble des réseaux (sans sauts) définis dans 1.2.7, dont les éléments seront appelés simplement réseaux
- PN_2^- : le sous-ensemble de PN_2^+ dont les éléments sont les réseaux qui ne contiennent aucun lien $\&$ et aucun lien $\oplus$.

A tout réseau avec sauts est associé un réseau (il suffit d'oublier les sauts), de telle sorte que tout réseau est en fait une classe d'équivalence de réseaux avec sauts. Toutes les e.e.r. de la définition 2.2.3, sauf l'étape $[ccad]$, sont parfaitement bien définies sur les réseaux (cf. aussi théorème 2.3.2). On peut donc définir dans PN_2^+ la règle de réécriture $\xrightarrow{\neg[ccad]}$, obtenue à partir de la réécriture usuelle en interdisant toute application d'une e.e.r. $[ccad]$. On notera $\xrightarrow{\neg[ccad]}$ la cōture réflexive et transitive de $\xrightarrow{\neg[ccad]}$. On dira aussi que σ est une suite de réductions $\neg[ccad]$ lorsque $\forall [x] \in \sigma, [x] \neq [ccad]$.

Soient (R, j_R) et $(R', j_{R'})$ deux réseaux avec sauts t.q. $(R, j_R) \xrightarrow{\sigma} (R', j_{R'})$, où σ est une suite de réductions $\neg[ccad]$. On a alors pour les réseaux $R, R' : R \xrightarrow{\sigma} R'$. Puisqu'en l'absence d'e.e.r. $[ccad]$ la notion de saut n'a aucune influence sur le processus d'élimination des coupures, la réciproque aussi est vraie : si (R, j_R) est un réseau avec sauts et R' est un réseau t.q. $R \xrightarrow{\sigma} R'$ (où σ est une suite de réductions $\neg[ccad]$), alors il existe une fonction de saut $j_{R'}$ t.q. $(R', j_{R'})$ est un réseau avec sauts et $(R, j_R) \xrightarrow{\sigma} (R', j_{R'})$ (il suffira de suivre la définition 2.2.8, et d'utiliser le théorème 2.3.5).

Nous appliquerons très souvent le lemme de substitution (lemme 3.1.1) : lorsque une e.e.r. $[\forall/\exists]$ est concernée par une de ces applications, alors les formules de certains réseaux doivent être considérées à des substitutions près.

Nous allons utiliser les résultats suivants :

4.0.5. THÉORÈME. ([Girard 87], [Danos 90]). Les réseaux de PN_2^- sont fortement normalisables (c.à.d. que toute suite d'e.e.r. issue de l'un d'entre eux est finie).

Démonstration : C'est essentiellement la preuve de normalisation forte de Girard pour le système F, adaptée aux réseaux ([Girard 87]). Mais la syntaxe linéaire rend indispensable pour mener à bien la preuve un résultat supplémentaire, connu sous le nom de "théorème de standardisation" ou "théorème de striction". Pour le fragment PN_2^- , ce résultat a été prouvé par Vincent Danos (cf. [Danos 90]). $\square$

4.0.6. THÉORÈME. ([Bechet 94]). Il existe une traduction des réseaux de

PN_2^+ dans les réseaux de PN_2^- , notée $[\cdot]$, et telle que : si R est un réseau de PN_2^+ et $R \xrightarrow{[x]} R'$, avec $[x] \neq [ccad]$, alors on a dans PN_2^- $[R] \xrightarrow{n} [R']$, où $n > 0$.

4.0.7. COROLLAIRE. Les réseaux de PN_2^+ sont fortement normalisables, relativement à la règle de réécriture $\xrightarrow{\neg[ccad]}$.

4.1 La réduction $\neg[ccad]$

Le résultat de Bechet étant formulé en calcul des séquents, et n'étant pas à ce jour publié, nous commencerons par le redémontrer rapidement.

Il repose sur le codage des connecteurs additifs de PN_2^+ dans PN_2^- que voici :

$$\begin{aligned} U \oplus' V &= \forall X.?(X \otimes U) \wp?(X \otimes V) \wp X^\perp \\ U \&' V &= \exists X.!(X^\perp \wp U) \otimes !(X^\perp \wp V) \otimes X. \end{aligned}$$

4.1.1. PROPOSITION. Les formules suivantes sont prouvables (dans PN_2^+) :

- (i) $(U \oplus V)^\perp \wp (U \oplus' V)$ et $(U \oplus' V)^\perp \wp (U \oplus V)$
- (ii) $(U \& V)^\perp \wp (U \&' V)$ et $(U \&' V)^\perp \wp (U \& V)$.

Démonstration : Laisée au lecteur (ou bien cf. [Bechet 94]). □

Pour toute formule U , on va définir la formule $[U]$, obtenue à partir de U en substituant $\&'$ à $\&$ et $\oplus'$ à $\oplus$. $[U]$ sera alors une formule qui ne contiendra aucune occurrence de connecteur additif.

4.1.2. DÉFINITION. Si U est une formule, la formule $[U]$ se définit par induction :

- $[X] = X$ si X est une variable propositionnelle.
- $[U \otimes V] = [U] \otimes [V]$ et $[U \wp V] = [U] \wp [V]$ pour les multiplicatifs.
- $[!U] = ![U]$ et $[?U] = ?[U]$ pour les exponentielles.
- $[\forall X.U] = \forall X.[U]$ et $[\exists X.U] = \exists X.[U]$ pour les quantificateurs.
- $[U \oplus V] = [U] \oplus' [V] = \forall X.?(X \otimes U) \wp?(X \otimes V) \wp X^\perp$ avec X non libre dans U ni dans V .
- $[U \& V] = [U] \&' [V] = \exists X.!(X^\perp \wp U) \otimes !(X^\perp \wp V) \otimes X$ avec X non libre dans U ni dans V .

On déduit de la proposition précédente qu'une formule A est prouvable dans PN_2^+ ssi $[A]$ est prouvable dans PN_2^- .

Dans la suite, nous écrirons de temps en temps des preuves du calcul des séquents au lieu de réseaux, mais il convient de considérer celles-ci comme une “notation” pour les réseaux de PN_2^+ : ce n’est qu’une façon d’améliorer la lisibilité. Dans la définition suivante, on utilise le théorème de séquentialisation pour PN_2^+ (le théorème 1.2.8). Si $\Gamma = A_1, \dots, A_n$ est un multiensemble de formules, alors on note $[\Gamma]$ le multiensemble $[A_1], \dots, [A_n]$.

4.1.3. DÉFINITION. Soit R un réseau de PN_2^+ de conclusions $U_1, \dots, U_n$. Le réseau $[R]$ de PN_2^- se définit par induction sur une séquentialisation π de R :

- si la dernière règle de π est une règle $\oplus$ (par exemple une règle $\oplus_1$) :

$$\frac{\begin{array}{c} R_1 \\ \vdots \\ \vdash \Gamma, U \end{array}}{\vdash \Gamma, U \oplus V} \oplus_1$$

alors on définit $[R]$ comme :

$$\frac{\begin{array}{c} [R_1] \\ \vdots \\ \frac{\vdash X, X^\perp \quad \vdash [\Gamma], [U]}{\vdash [\Gamma], X \otimes [U], X^\perp} \otimes \\ \frac{\vdash [\Gamma], ?(X \otimes [U]), ?(X \otimes [V]), X^\perp}{\vdash [\Gamma], ?(X \otimes [U]) \wp ?(X \otimes [V]) \wp X^\perp} ? \\ \frac{\vdash [\Gamma], ?(X \otimes [U]) \wp ?(X \otimes [V]) \wp X^\perp}{\vdash [\Gamma], \forall X. ?(X \otimes [U]) \wp ?(X \otimes [V]) \wp X^\perp} \wp \\ \hline \vdash [\Gamma], [U \oplus V] \quad \mathbf{Def} \end{array}}$$

- si R est une boîte additive (c.à.d. que la dernière règle de π est une règle $\&$) :

$$\frac{\begin{array}{c} R_1 \\ \vdots \\ \vdash \Gamma, U \end{array} \quad \begin{array}{c} R_2 \\ \vdots \\ \vdash \Gamma, V \end{array}}{\vdash \Gamma, U \& V} \&$$

alors on définit $[R]$ comme :

$$\begin{array}{c}
[R_1] \qquad \qquad \qquad [R_2] \\
\vdots \qquad \qquad \qquad \vdots \\
\frac{\vdash [\Gamma], [U]}{\vdash C, [U]} \wp \text{ (?W)} \qquad \frac{\vdash [\Gamma], [V]}{\vdash C, [V]} \wp \text{ (?W)} \qquad \eta \\
\frac{\vdash C, [U]}{\vdash C \wp [U]} \wp \qquad \frac{\vdash C, [V]}{\vdash C \wp [V]} \wp \\
\frac{\vdash C \wp [U]}{\vdash!(C \wp [U])} ! \qquad \frac{\vdash C \wp [V]}{\vdash!(C \wp [V])} ! \qquad \vdash [\Gamma], C^\perp \\
\hline
\frac{\vdash [\Gamma], !(C \wp [U]) \otimes !(C \wp [V]) \otimes C^\perp}{\vdash [\Gamma], \exists X.!(X^\perp \wp [U]) \otimes !(X^\perp \wp [V]) \otimes X} \otimes \\
\hline
\frac{\vdash [\Gamma], \exists X.!(X^\perp \wp [U]) \otimes !(X^\perp \wp [V]) \otimes X}{\vdash [\Gamma], [U \& V]} \mathbf{Def}
\end{array}$$

où lorsque $\Gamma = A_1, \dots, A_h$ on prend comme formule C la formule $[A_1] \wp \dots \wp [A_h]$ et comme preuve η la preuve obtenue à partir des axiomes $\vdash [A_i], [A_i]^\perp$ en appliquant $n - 1$ règles $\otimes$, tandis que lorsque $\Gamma = \emptyset$ on prend comme formule C une formule $?D$ quelconque t.q. il existe une preuve η de $!D^\perp = C^\perp$ (prendre par exemple $D^\perp = E \wp E^\perp$ où E est une formule).

- si la dernière règle de π est toute autre règle L ($0 \leq k \leq 2$) :

$$\begin{array}{c}
R_1 \qquad \qquad \qquad R_k \\
\vdots \qquad \qquad \qquad \vdots \\
\vdash \Gamma_1 \qquad \dots \qquad \vdash \Gamma_k \\
\hline
\vdash U_1, \dots, U_n \quad L
\end{array}$$

alors on définit $[R]$ comme :

$$\begin{array}{c}
[R_1] \qquad \qquad \qquad [R_k] \\
\vdots \qquad \qquad \qquad \vdots \\
\vdash [\Gamma_1] \qquad \dots \qquad \vdash [\Gamma_k] \\
\hline
\vdash [U_1], \dots, [U_n] \quad L
\end{array}$$

Le lemme suivant est le lemme crucial pour prouver la forte normalisation de $\neg[ccad]$:

4.1.4. LEMME. Si R est un réseau de PN_2^+ et $R \xrightarrow{[x]} R'$ (avec bien sûr $[x] \neq [ccad]$), alors on a dans PN_2^- $[R] \xrightarrow{n} [R']$, où $n > 0$.

Démonstration : Remarquons tout d'abord que le lien coupure x de R apparaît dans $[R]$: appelons y ce lien coupure de $[R]$. Soit n_1 et n_2 (resp. m_1 et m_2) les deux liens dont les conclusions sont les deux prémisses de x dans R (resp. de y dans $[R]$). Des définitions 4.1.2 et 4.1.3, on déduit que m_i est en fait n_i , sauf dans les trois cas suivants :

- (1) n_i est un lien $\&$
- (2) n_i est un lien $\oplus$
- (3) n_i est un lien *ccad*.

Si ni n_1 ni n_2 ne satisfait une des trois conditions ci-dessus, alors on a $[R] \xrightarrow{[y]} [R']$.

Si (1), (2), ou (3) est satisfaite par un des liens n_1, n_2 , alors $[x]$ est nécessairement une e.e.r. axiome, $[\&/\oplus]$, ou $[ccad]$. Ce dernier cas est à exclure ($[x] \neq [ccad]$).

Si $[x]$ est une e.e.r. axiome et (par exemple) n_2 est un lien axiome, m_2 est aussi un lien axiome et on a $[R] \xrightarrow{[ax]} [R']$. La seule chose qui reste à vérifier est que le lemme est vrai lorsque $[x] = [\&/\oplus]$. Si $[x] = [\&/\oplus_1]$, on a :

$$\begin{array}{c}
\begin{array}{c} [R_1] \\ \vdots \\ \frac{\vdash [\Gamma], [U]}{\vdash C, [U]} \wp, (?W) \\ \frac{\vdash C, [U]}{\vdash C \wp [U]} \wp \\ \frac{\vdash C \wp [U]}{\vdash (C \wp [U])} ! \end{array} \quad
\begin{array}{c} [R_2] \\ \vdots \\ \frac{\vdash [\Gamma], [V]}{\vdash C, [V]} \wp, (?W) \\ \frac{\vdash C, [V]}{\vdash C \wp [V]} \wp \\ \frac{\vdash C \wp [V]}{\vdash (C \wp [V])} ! \end{array} \quad
\eta \\
\frac{\vdash [\Gamma], (C \wp [U]) \otimes (C \wp [V]) \otimes C^\perp}{\vdash [\Gamma], \exists X. (X^\perp \wp [U]) \otimes (X^\perp \wp [V]) \otimes X} \otimes \quad
\frac{\vdash [\Gamma], (C \wp [U]) \otimes (C \wp [V]) \otimes C^\perp}{\vdash [\Gamma], \exists X. (X^\perp \wp [U]) \otimes (X^\perp \wp [V]) \otimes X} \exists X = C^\perp \\
\frac{\vdash [\Gamma], \exists X. (X^\perp \wp [U]) \otimes (X^\perp \wp [V]) \otimes X}{\vdash [\Gamma], [U \& V]} \mathbf{Def}
\end{array}
\quad
\begin{array}{c} [R_3] \\ \vdots \\ \frac{\vdash X, X^\perp}{\vdash [\Delta], [U^\perp]} \otimes \\ \frac{\vdash [\Delta], X \otimes [U^\perp], X^\perp}{\vdash [\Delta], ?(X \otimes [U^\perp]), ?(X \otimes [V^\perp]), X^\perp} ? \\ \frac{\vdash [\Delta], ?(X \otimes [U^\perp]), ?(X \otimes [V^\perp]), X^\perp}{\vdash [\Delta], ?(X \otimes [U^\perp]) \wp ?(X \otimes [V^\perp]) \wp X^\perp} \wp \\ \frac{\vdash [\Delta], ?(X \otimes [U^\perp]) \wp ?(X \otimes [V^\perp]) \wp X^\perp}{\vdash [\Delta], \forall X. ?(X \otimes [U^\perp]) \wp ?(X \otimes [V^\perp]) \wp X^\perp} \vee \\ \frac{\vdash [\Delta], \forall X. ?(X \otimes [U^\perp]) \wp ?(X \otimes [V^\perp]) \wp X^\perp}{\vdash [\Delta], [U^\perp \oplus V^\perp]} \mathbf{Def}
\end{array}
\quad
\frac{\vdash [\Gamma], [U \& V] \quad \vdash [\Delta], [U^\perp \oplus V^\perp]}{\vdash [\Gamma], [\Delta]} \mathbf{cut}
\end{array}$$

qui se réduit en plusieurs étapes en

$$\begin{array}{c}
[R_1] \quad [R_3] \\ \vdots \quad \vdots \\ \frac{\vdash [\Gamma], [U] \quad \vdash [\Delta], [U^\perp]}{\vdash [\Gamma], [\Delta]} \mathbf{cut}
\end{array}$$

□

4.1.5. PROPOSITION. Les réseaux PN_2^+ sont fortement normalisables, relativement à la règle de réécriture $\xrightarrow{\neg[ccad]}$.

Démonstration : C'est une conséquence immédiate du lemme précédent : une suite infinie d'e.e.r. $\neg[ccad]$ dans PN_2^+ produirait une suite infinie d'e.e.r. dans PN_2^- , ce qui contredit le théorème 4.0.5. □

4.1.6. REMARQUE. Soit R un réseau et B une boîte additive de R . Supposons que $R \xrightarrow{\sigma} R'$, où σ est une suite de réductions $\neg[ccad]$. Aucun résidu B' de B dans R' ne peut contenir comme sous-réseau un résidu B'' de B

différent de B' : tout simplement parce que, en l'absence de $[ccad]$, il n'y a aucun moyen d'entrer dans une boîte additive.

Soit (R, j_R) un réseau avec sauts et supposons que $R \xrightarrow{[x]} R'$, où $[x] \neq [ccad]$. A cause des différentes façons possibles de définir la fonction $j_{R'}$ associée à l'e.e.r. $[x]$ (suivant la définition 2.2.8), la proposition suivante ne pourra pas être vraie dans PN_2 .

4.1.7. PROPOSITION. Dans PN_2^+ , la règle de réécriture $\neg[ccad] \rightarrow$ est confluente.

Démonstration : C'est une conséquence de la forte normalisation et de la confluence locale de $\neg[ccad] \rightarrow$, où la confluence locale est la propriété suivante des réseaux : si R est un réseau t.q. $R \xrightarrow{[x_1]} R_1$ et $R \xrightarrow{[x_2]} R_2$, alors il existe deux suites de réductions σ_1 et σ_2 et un réseau R' t.q. $R_1 \xrightarrow{\sigma_1} R'$ et $R_2 \xrightarrow{\sigma_2} R'$. Ceci peut se démontrer en examinant tous les cas possibles (sans surprises). $\square$

4.1.8. PROPOSITION. Soit R un réseau sans occurrences du connecteur $\&$ ni du quantificateur $\exists$ dans ses conclusions.

Si R est normal pour la règle de réécriture $\neg[ccad] \rightarrow$, (c.à.d. qu'il n'y a dans R que des liens coupure de type $ccad$), alors R est sans coupures.

Démonstration : C'est plus ou moins évident, et J.-Y. Girard l'avait déjà remarqué dans [Girard 95a] : R ne peut pas contenir seulement des liens coupure de type $ccad$. En effet, si c_1 est un lien coupure dont une prémisses est une porte coad d'une boîte additive B_1 , alors il existe un chemin Φ_1 ayant la porte principale de B_1 comme premier lien. Ce chemin doit traverser un lien coupure c_2 (se souvenir qu'il n'y a pas d'occurrences du connecteur $\&$ ni du quantificateur $\exists$ dans les conclusions de R). Mais alors, c_2 est un lien coupure dont une prémisses est une porte coad d'une boîte additive B_2 , et il existe un chemin Φ_2 ayant la porte principale de B_2 comme premier lien, et qui doit traverser un lien coupure $c_3 \dots$ bien sûr ceci ne peut pas continuer éternellement (à cause de l'acyclicité de tous les graphes de correction d'un réseau, cf. définition 1.2.7). $\square$

4.2 La réduction paresseuse

Nous allons à présent faire une petite digression sur la stratégie dite "paresseuse" (*lazy* en anglais) d'élimination des coupures, à cause de son impor-

tance dans les récents travaux sur les systèmes logiques à “basse complexité” ELL et LLL (cf. notamment [Girard 95a]).

4.2.1. DÉFINITION. Soit x un lien coupure du réseau R . Nous dirons que x est un **lien coupure paresseux** lorsque x est à profondeur additive zéro dans R , et x n’est pas *seulement* de type $ccad$ (se souvenir de la remarque 2.2.2 : un lien coupure peut être à la fois de type axiome et de type $ccad$). De la même façon, si x est un lien coupure paresseux, alors nous dirons que $[x]$ est une **e.e.r. paresseuse**, *lorsque* $[x] \neq [ccad]$ (ceci devant être précisé, toujours à cause de la remarque 2.2.2). Nous dirons d’une suite de réductions σ que c’est une **suite de réductions paresseuses** lorsque chaque e.e.r. de σ est paresseuse. On notera $\longrightarrow_l$ cette règle de réécriture, et $\rightarrow_l$ sa clôture réflexive et transitive.

4.2.2. PROPOSITION. Soit R un réseau ne contenant aucune occurrence du connecteur $\&$ ni du quantificateur $\exists$ dans ses conclusions. Alors :

- (1) Toute suite d’e.e.r. paresseuses issue de R est finie.
- (2) Toute forme normale paresseuse de R est sans coupures : si un réseau T est normal pour $\longrightarrow_l$, (c.à.d. qu’il ne contient que des liens coupures de type $ccad$ ou de profondeur additive supérieure à zéro dans T), alors T est sans coupures.
- (3) Il existe une unique forme normale paresseuse de R (qui est sans coupures).

Démonstration : (1) : Puisque toute e.e.r. paresseuse est aussi une e.e.r. $\neg[ccad]$, la règle de réécriture $\longrightarrow_l$ est fortement normalisante (conséquence de la proposition 4.1.5).

(2) : Identique à la preuve de 4.1.8.

(3) : Soient R_0 et R_1 deux formes normales paresseuses de R . (2) implique que R_0 et R_1 sont sans coupures. Puisque une suite de réductions paresseuses est aussi une suite de réductions $\neg[ccad]$, R_0 et R_1 sont aussi des formes normales pour la règle de réécriture $\neg[ccad]$, qui est confluente. Donc $R_0 = R_1$. $\square$

4.2.3. REMARQUE. La procédure paresseuse d’élimination des coupures est “presque” celle qui est définie dans [Girard 95b] pour éliminer les coupures dans le système LLL , et elle coïncide exactement avec la procédure définie par Girard dans le cas de $MALL_2$ (le fragment de PN_2 qui ne fait usage que des connecteurs additifs et des quantificateurs du second ordre). Dans

[Girard 95a], on montre que (en faisant un peu attention aux substitutions dans les e.e.r. $[\forall/\exists]$) pour $MALL_2$ cette procédure termine en temps linéaire. Le “défaut” de la procédure définie par Girard est qu’elle ne termine pas toujours sur un réseau sans coupures. Mais dans le cas où les conclusions du réseau qu’on normalise ne contiennent aucune occurrence du connecteur $\&$ ni du quantificateur $\exists$, (2) de la proposition 4.2.2 nous dit précisément que la procédure paresseuse termine sur un réseau sans coupures. Ce cas est en fait très important car les données (telles que par exemple les entiers) sont le plus souvent typées par des formules qui ne contiennent aucune occurrence du connecteur $\&$ ni du quantificateur $\exists$ (cf. aussi [DJ 99]).

La remarque suivante est une conséquence immédiate de la proposition précédente et de la confluence de $\neg[ccad]$ (proposition 4.1.7).

4.2.4. REMARQUE. Soit R un réseau ne contenant aucune occurrence du connecteur $\&$ ni du quantificateur $\exists$ dans ses conclusions. L’unique forme normale $\neg[ccad]$ de R peut être atteinte de façon paresseuse.

Nous montrerons (théorème 4.3.5) que l’unique forme normale $\neg[ccad]$ de R est en fait aussi l’unique forme normale du réseau avec sauts (R, j) (où j est une fonction de saut quelconque), et ce que nous venons de dire c’est qu’on peut l’atteindre de façon paresseuse.

4.3 Le théorème de confluence faible

4.3.1. LEMME. Soit (R, j_R) un réseau avec sauts t.q. $(R, j_R) \xrightarrow{[ccad]} (R_1, j_{R_1})$. Soit $\mathbb{B}$ la boîte additive de R attractive pour l’e.e.r. $[ccad]$, et (S, j_S) le sous-réseau de (R, j_R) avalé par $\mathbb{B}$. Appelons (α, j_α) le sous-réseau de (R, j_R) obtenu en connectant $\mathbb{B}$ et S par le lien coupure qui réduit l’e.e.r. $[ccad]$, et B l’unique résidu de $\mathbb{B}$ dans (R_1, j_{R_1}) . On a évidemment $R = R_1[\alpha/B]$.

Supposons que $R_1 \xrightarrow{\sigma} T$, où σ est une suite de réductions $\neg[ccad]$ qui ne contient aucune e.e.r. réduisant une coupure à l’intérieur d’un résidu de B . Si on appelle $B_1, \dots, B_n$ les n résidus de B dans T , alors $R \xrightarrow{\sigma} T[\alpha/B_1, \dots, \alpha/B_n]$.

Démonstration : Essentiellement, le lemme est vrai parce que lors d’une suite de réductions $\neg[ccad]$, un résidu de B peut seulement être dupliqué (par une étape $[co]$), effacé (par une étape $[w]$ ou $[\&/\oplus]$) ou bien ouvert.

Remarquons aussi que la seule raison pour laquelle on ne peut pas appliquer directement le lemme de substitution (lemma 3.1.1) est que parmi les

e.e.r. de σ il pourrait y avoir $[x] = [\&/\oplus]$ qui ouvre un résidu de B . Ce qu'on va montrer, c'est précisément que cette éventualité ne présente pas d'inconvénients : $[x]$ a le même effet lorsqu'on substitue α à ce résidu de B . On prouve le lemme par induction sur σ .

Si $\sigma = \emptyset$, alors le résultat est évident.

Autrement, supposons que $R_1 \xrightarrow{\sigma'} T' \xrightarrow{[x]} T$, et soient $B_1, \dots, B_n$ les résidus de B dans T' . Par hypothèse d'induction $R \xrightarrow{\sigma'} T'[\alpha/B_1, \dots, \alpha/B_n]$.

Il faut montrer que $\forall i \in \{1, \dots, n\}$, $T'[\alpha/B_i] \xrightarrow{[x]} T[\alpha/B_i^1, \dots, \alpha/B_i^{h_i}]$, où $B_i^1, \dots, B_i^{h_i}$ sont les h_i résidus de B_i dans T ($h_i \in \{0, 1, 2\}$).

Si x n'est pas un lien coupure $\&/\oplus$ dont une prémisse est la porte $\&$ d'un résidu de B , alors on applique tout simplement le lemme de substitution à chaque B_i .

Supposons donc que x soit un lien coupure $\&/\oplus$ dont une prémisse est la porte $\&$ de B_k . Remarquons d'abord qu'on peut toujours appliquer le lemme de substitution à tous les B_j , pour $j \neq k$ (parce que aucune conclusion de B_j ne peut être conclusion d'un lien $\oplus$, et bien sûr aussi à cause de la remarque 4.1.6). On vérifie alors simplement que $T'[\alpha/B_k] \xrightarrow{[\&/\oplus]} T$. Puisqu'il n'y a pas de résidus de B_k dans T , ceci permet de conclure. $\square$

4.3.2. LEMME. Soit B une boîte additive du réseau R . Supposons que $R \xrightarrow{\sigma_1} R_1 \xrightarrow{[x]} R_2 \xrightarrow{\sigma_2} R_0$ soit une suite de réductions $\neg[ccad]$, que R_0 ne contienne aucun résidu de B , et que σ_2 ne contienne aucune e.e.r. réduisant une coupure à l'intérieur d'un résidu de B .

Alors, il existe une suite de réductions $\neg[ccad] \sigma$ t.q. $R \xrightarrow{\sigma_1} R_1 \xrightarrow{\sigma} R_0$, et σ ne contient aucune e.e.r. réduisant une coupure dans un résidu de B .

Démonstration : Il n'y a quelque chose à prouver que lorsque le lien coupure x est contenu dans le résidu B_1 de B dans R_1 . Soit B_2 l'unique résidu de B_1 idans R_2 .

L'idée est que, puisque finalement tous les résidus de B_2 disparaîtront, on a pour chacun de ces résidus une des deux possibilités suivantes :

- le résidu est effacé par une e.e.r. de σ_2 , auquel cas $[x]$ est inutile (pour ce résidu)
- le résidu est ouvert par une e.e.r. $[\&/\oplus]$, auquel cas, soit $[x]$ est de nouveau inutile, soit "c'est pareil" d'appliquer $[x]$ au début de la suite de réductions ou immédiatement après l'e.e.r. $[\&/\oplus]$.

Plus formellement, on prouve que pour toute décomposition de σ_2

$R_2 \xrightarrow{\tau} T \xrightarrow{\sigma_2 \setminus \tau} R_0$, il existe une suite de réductions $\neg[ccad] \epsilon$ t.q.

$R_1 = R_2[B_1/B_2] \xrightarrow{\epsilon} T[B_1/B_2^1, \dots, B_1/B_2^n]$, et $\forall [z] \in \epsilon$, z n'est contenu dans aucun résidu de B (nous sommes en train de parler des résidus de la boîte B de R), où $B_2^1, \dots, B_2^n$ sont les n résidus de B_2 dans T (relativement à la suite de réductions τ). La conclusion sera alors une conséquence immédiate du cas $\tau = \sigma_2$.

On procède par induction sur τ .

Si $\tau = \emptyset$, alors $T = R_2$ et le résultat est immédiat avec $\epsilon = \tau = \emptyset$.

Autrement, $R_2 \xrightarrow{\tau_1} T_1 \xrightarrow{[y]} T$. Soit $B_2^1, \dots, B_2^m$ les résidus de B_2 dans T_1 . Par hypothèse d'induction $R_1 = R_2[B_1/B_2] \xrightarrow{\epsilon_1} T_1[B_1/B_2^1, \dots, B_1/B_2^m]$ où $\forall [z] \in \epsilon_1$, z n'est contenu dans aucun résidu de B .

$\forall i \in \{1, \dots, m\}$, soit $B_2^{1,i}, \dots, B_2^{h_i,i}$ les h_i résidus de B_2^i dans T ($h_i \in \{0, 1, 2\}$). Si $[y]$ n'ouvre aucune des boîtes $B_2^1, \dots, B_2^m$, alors le lemme de substitution permet d'affirmer que, $\forall i \in \{1, \dots, m\}$, $T_1[B_1/B_2^i] \xrightarrow{[y]} T[B_1/B_2^{1,i}, \dots, B_1/B_2^{h_i,i}]$, et on prendra $\epsilon = [y] \circ \epsilon_1$.

Supposons maintenant que $[y] = [\&/\oplus]$ ouvre la boîte B_2^k . Alors $\forall i \in \{1, \dots, m\}$, $i \neq k$, $h_i = 1$; et il n'existe aucun résidu de B_2^k dans T . Pour $i \neq k$ le lemme de substitution peut toujours s'appliquer, et on a $T_1[B_1/B_2^i] \xrightarrow{[y]} T[B_1/B_2^i]$. Pour $i = k$, puisque dans σ_2 il n'y a aucune e.e.r. réduisant une coupure dans une composante d'un résidu de B_2 , de $B_1 \xrightarrow{[x]} B_2$, on déduit $B_1 \xrightarrow{[x]} B_2^k$. Soit T_2 t.q. $T_1[B_1/B_2^k] \xrightarrow{[y]} T_2$. On a les deux possibilités suivantes :

- si x est un lien de la composante de B_1 effacée par $[y] = [\&/\oplus]$, alors $T_2 = T$
- si x est un lien de la composante de B_1 qui n'est pas effacée par $[y] = [\&/\oplus]$, alors $T_2 \xrightarrow{[x]} T$. (Remarquons que dans ce cas x n'est contenu dans aucun résidu de B dans T_2 , car autrement dans T_1 un résidu de B serait contenu dans B_2^k , en contredisant de la sorte la remarque 4.1.6).

Dans le premier cas on prendra $\epsilon = [y] \circ \epsilon_1$, tandis que dans le deuxième on prendra $\epsilon = [x] \circ [y] \circ \epsilon_1$.

Pour être vraiment précis, il faudrait aussi vérifier que y n'est contenu dans aucun résidu de B relativement à $\epsilon_1 \circ \sigma_1$: ceci est clairement une conséquence de la construction de ϵ_1 . $\square$

4.3.3. LEMME. Soit R un réseau et B une boîte additive de R . Si $R \xrightarrow{\sigma'} R_0$, où σ' est une suite de réductions $\neg[ccad]$ et il n'y a aucun résidu de B dans R_0 , alors il existe une suite de réductions $\neg[ccad]$ σ t.q. $R \xrightarrow{\sigma} R_0$ et σ ne contient aucune e.e.r. réduisant un lien coupure contenu dans un résidu de

B .

Démonstration : Soit $n(\sigma')$ le nombre d'e.e.r. $[y]$ de σ' t.q. y appartient à un résidu de B . On montre le lemme par induction sur $n(\sigma')$.

Si $n(\sigma') = 0$, alors on prendra $\sigma = \sigma'$.

Autrement, soit $[x]$ la dernière e.e.r. de σ' t.q. x est contenu dans un résidu de B . On a la décomposition suivante de $\sigma' : R \xrightarrow{\sigma_1} R_1 \xrightarrow{[x]} R_2 \xrightarrow{\sigma_2} R_0$, où σ_2 ne contient aucune e.e.r. réduisant un lien coupure à l'intérieur d'un résidu de B . Le lemme précédent nous fournit une suite de réductions $\neg[ccad]$ τ t.q. $R \xrightarrow{\sigma_1} R_1 \xrightarrow{\tau} R_0$ et τ ne contient aucune e.e.r. réduisant une coupure dans un résidu de B . Alors, si on définit $\epsilon := \tau \circ \sigma_1$, on obtient $R \xrightarrow{\epsilon} R_0$ et $n(\epsilon) < n(\sigma')$. On peut alors conclure en appliquant l'hypothèse d'induction à ϵ . $\square$

4.3.4. LEMME. Soit R un réseau ne contenant aucune occurrence du connecteur $\&$ ni du quantificateur $\exists$ dans ses conclusions, et soit (R, j_R) un réseau avec sauts.

Si $(R, j_R) \xrightarrow{[ccad]} (R_1, j_{R_1})$, alors il existe une suite de réductions $\neg[ccad]$ σ , et un réseau sans coupures R_0 t.q. $R \xrightarrow{\sigma} R_0$ et $R_1 \xrightarrow{\sigma} R_0$.

Démonstration : Soit $\mathbb{B}$ la boîte additive de R attractive pour l'e.e.r. $[ccad]$, et (S, j_S) le sous-réseau (avec sauts) de (R, j_R) avalé par $\mathbb{B}$. Appelons (α, j_α) le sous-réseau (avec sauts) de R obtenu en connectant $\mathbb{B}$ et S par le lien coupure que $[ccad]$ réduit, et B l'unique résidu de $\mathbb{B}$ dans (R_1, j_{R_1}) . On a $R_1[\alpha/B] = R$.

Soit σ' une suite de réductions $\neg[ccad]$ issue de R_1 et conduisant à l'unique forme normale $\neg[ccad]$ R_0 de R_1 . Nous savons par la proposition 4.1.8 que R_0 est sans coupures. La propriété de la sous-formule nous garantit qu'il n'y a pas de résidu de B (relativement à σ') dans R_0 . On peut alors appliquer le lemme 4.3.3 : il existe une suite de réductions $\neg[ccad]$ σ t.q. $R_1 \xrightarrow{\sigma} R_0$ et σ ne contient aucune e.e.r. réduisant une coupure dans un résidu de B . On peut maintenant appliquer le lemme 4.3.1 à σ et affirmer que $R \xrightarrow{\sigma} R_0$ (rappelons qu'il n'y a pas de résidu de B dans R_0). $\square$

On va maintenant prouver le théorème principal du chapitre.

4.3.5. THÉORÈME. Soit R un réseau ne contenant aucune occurrence du connecteur $\&$, ni du quantificateur $\exists$ dans ses conclusions.

Soient j_1 et j_2 deux fonctions t.q. (R, j_1) et (R, j_2) soient deux réseaux avec sauts.

Si $(R, j_1) \xrightarrow{\sigma_1} (R_0, j_1^0)$ et $(R, j_2) \xrightarrow{\sigma_2} (T_0, j_2^0)$, où R_0 et T_0 sont des réseaux sans coupures, alors $R_0 = T_0$.

En particulier (prendre $j_1 = j_2$), si $(R, j_1) \xrightarrow{\sigma_1} (R_0, j_1^0)$ et $(R, j_1) \xrightarrow{\sigma_2} (T_0, j_2^0)$, alors $R_0 = T_0$.

Démonstration : On montre que si (R_0, j_{R_0}) est une forme normale de (R, j_R) , alors il existe une suite de réductions $\neg[ccad]$ τ , t.q. $R \xrightarrow{\tau} R_0$. La confluence de $\xrightarrow{\neg[ccad]}$ (proposition 4.1.7) sera alors suffisante pour conclure : si $(R, j_1) \xrightarrow{\sigma_1} (R_0, j_1^0)$ et $(R, j_2) \xrightarrow{\sigma_2} (T_0, j_2^0)$, alors on aura $R \xrightarrow{\tau_1} R_0$ et $R \xrightarrow{\tau_2} T_0$ (où τ_1 et τ_2 sont des suites de réductions $\neg[ccad]$), donc $R_0 = T_0$.

Soit τ' une suite de réductions t.q. $(R, j_R) \xrightarrow{\tau'} (R_0, j_{R_0})$. On montre l'existence de τ par induction sur $l(\tau') :=$ nombre d'e.e.r. de τ' .

Si $l(\tau') = 0$, alors on prendra $\tau = \tau' = \emptyset$.

Autrement $(R, j_R) \xrightarrow{[x]} (R_1, j_{R_1}) \xrightarrow{\tau'_1} (R_0, j_{R_0})$. Par hypothèse d'induction (bien sûr (R_1, j_{R_1}) est un réseau avec sauts ne contenant aucune occurrence du connecteur $\&$, ni du quantificateur $\exists$ dans ses conclusions) il existe une suite de réductions $\neg[ccad]$ τ_1 t.q. $R_1 \xrightarrow{\tau_1} R_0$. Si $[x] \neq [ccad]$, alors on a fini. Si $[x] = [ccad]$, alors on applique le lemme 4.3.4 : il existe une suite de réductions $\neg[ccad]$ τ et un réseau sans coupures T_0 t.q. $R \xrightarrow{\tau} T_0$ et $R_1 \xrightarrow{\tau} T_0$.

La confluence de $\xrightarrow{\neg[ccad]}$, permet de déduire que $R_0 = T_0$ (on rappelle que R_0 et T_0 sont tous les deux sans coupures). $\square$

4.3.6. COROLLAIRE. (Confluence) Soit (R, j) un réseau avec sauts ne contenant aucune occurrence du connecteur $\&$ ni du quantificateur $\exists$ dans ses conclusions.

Si $(R, j) \xrightarrow{\sigma_1} (R_1, j_1)$ et $(R, j) \xrightarrow{\sigma_2} (R_2, j_2)$, alors il existe un réseau sans coupures R_0 , deux fonctions de saut j_1^0 et j_2^0 , et deux suites de réductions τ_1 et τ_2 t.q. $(R_1, j_1) \xrightarrow{\tau_1} (R_0, j_1^0)$ and $(R_2, j_2) \xrightarrow{\tau_2} (R_0, j_2^0)$.

Démonstration : La normalisation forte de $\neg[ccad]$ permet d'affirmer qu'en appliquant le nombre nécessaire d'e.e.r. à (R_1, j_1) et à (R_2, j_2) , on obtiendra deux suites de réductions $\neg[ccad]$ τ_1 et τ_2 et deux réseaux avec sauts et sans coupures (T_1^0, j_1^0) et (T_2^0, j_2^0) t.q. $(R_1, j_1) \xrightarrow{\tau_1} (T_1^0, j_1^0)$ et $(R_2, j_2) \xrightarrow{\tau_2} (T_2^0, j_2^0)$. Le théorème précédent affirme que $T_1^0 = T_2^0$. $\square$

4.3.7. REMARQUE. (i) Le lecteur n'aura pas manqué de remarquer que la raison pour laquelle on demande qu'aucune conclusion du réseau ne contienne comme étiquette une formule contenant une occurrence du quantificateur existentiel du second ordre, est simplement que celui-ci pourrait "cacher" des occurrences du connecteur $\&$.

(ii) Le théorème précédent affirme essentiellement que toute façon "correcte" d'éliminer les coupures dans un réseau R (avec des conclusions comme il faut) conduit toujours à la même forme normale. Plus précisément, la suite suivante d'opérations conduit toujours au même résultat (le même réseau) :

- prendre une fonction de saut j t.q. (R, j) est un réseau avec sauts
- éliminer les coupures jusqu'à atteindre un réseau sans coupures
- oublier la fonction de saut.

(iii) A moins de modifier la syntaxe des additifs (et nous tenterons de le faire dans le chapitre 5), le théorème 4.3.5 semble bien être ce que de mieux on puisse obtenir dans un cadre aussi général que celui dans lequel nous nous sommes placés.

4.3.8. REMARQUE. On aurait aussi pu envisager une approche sémantique à la preuve du théorème 4.3.5 : il aurait été suffisant de prouver que deux réseaux de PN_2 sans coupures et sans occurrences de $\&$ et de $\exists$ sont interprétés par deux cliques différentes (dans la sémantique cohérente ensembliste ou multiensembliste). Une première étape aurait été de prouver ce résultat pour deux réseaux sans coupures dans le fragment multiplicatif et exponentiel. Mais... surprise! Ceci est tout simplement faux, comme nous le montrerons dans la suite de la thèse.

Chapitre 5

Les multiboîtes additives

Dans les chapitres précédents, nous nous sommes occupés du calcul associé aux démonstrations de LL dans le cadre (très) général des réseaux de PN_2 . Le lecteur n'aura pas manqué de remarquer que les problèmes que nous avons rencontrés sont tous liés à la présence des liens $\&$ et (dans une moindre mesure) à celle des liens $?w$, la coexistence des deux nous ayant contraint à pas mal d'équilibrisme...

Nous nous sommes donnés beaucoup de mal pour définir correctement les e.e.r. *[ccad]* (dans le chapitre 2), pour montrer dans le chapitre 4 que, sous certaines conditions, ces étapes sont totalement inutiles. Même si les conditions en question sont remplies dans des cas très importants (voir aussi la remarque 4.2.4), une question incontournable se pose : et si elles ne le sont pas ?

Le problème qui se cache derrière cette question est celui des réseaux additifs. C'est un problème *très* compliqué, et à ce jour (après 13 ans de logique linéaire) seules quelques rares solutions (partielles) ont été proposées. Dans [Girard 95a], une amélioration de la syntaxe est proposée et un critère de correction faisant usage de "poids" est introduit. Ce même critère a été considérablement simplifié dans [Laurent 99], dans le cadre des réseaux polarisés (voir chapitre 15).

Notre approche à la question est plus inspirée par la recherche d'une meilleure notion de normalisation que par celle d'un critère de correction. Le problème étant, bien sûr, l'e.e.r. *[ccad]*.

On va s'éloigner du cadre (relativement bien établi) des réseaux de PN_2 , et il ne s'agira donc pas tellement d'avoir des résultats "les plus généraux possibles", mais plutôt de se placer dans un fragment significatif de LL (contenant tout de même les additifs!) et essayer de voir si, en modifiant la

syntaxe, on arrive à dire quelque chose de mieux sur la normalisation. On va donc se placer, dans tout le chapitre, dans le fragment multiplicatif et additif de LL.

Nous commençons par une analyse sémantique de la non confluence dans PN_2 , et nous suggérons donc au lecteur de se familiariser d’abord avec les notions introduites dans les deux premiers paragraphes du chapitre 6. Nous introduisons ensuite la notion de “multiboîte” (inspirée par la construction sémantique), et les modifications qu’elle suggère dans les étapes d’élimination des coupures. Nous prouvons un théorème de forte normalisation de cette variante de la normalisation (théorème 5.3.8), et nous définissons une procédure d’élimination des coupures conduisant à une unique forme normale (théorème 5.5.15).

Bien que partiels, ces résultats semblent indiquer qu’une gestion de la normalisation en présence d’additifs est possible, malgré l’incroyable complexité de la notion de “superposition” sous-jacente au connecteur $\&$.

5.1 De la non confluence aux multiboîtes

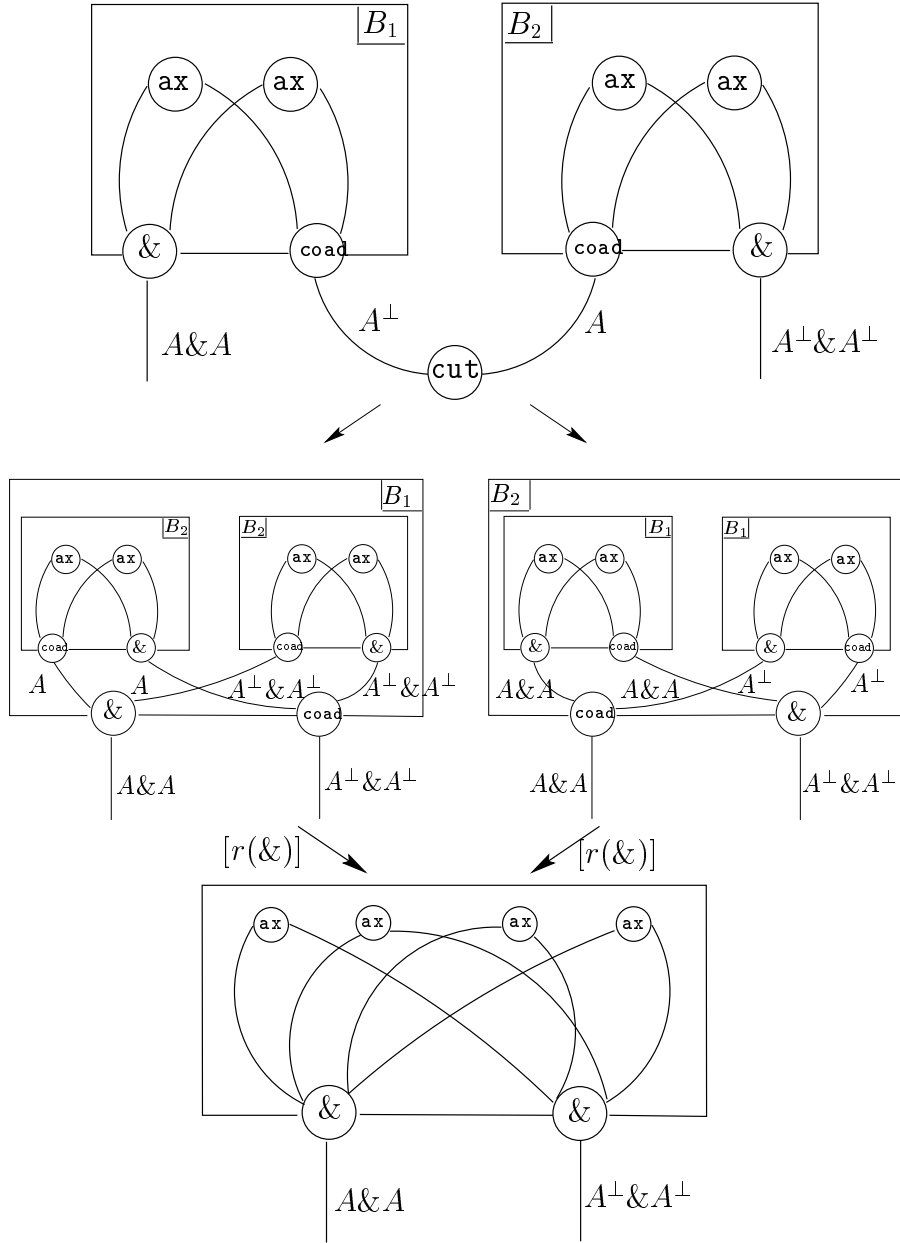
Dans tout le reste de la thèse, nous appellerons “type d’une arête a ” la formule associée à a , le terme “étiquette” ayant une autre signification (voir remarque 6.2.2). D’ailleurs, dans ce chapitre, nous confondrons bientôt une arête avec son type (voir 5.4.1).

Remarquons qu’en l’absence d’exponentielles, les sémantiques relationnelle et cohérente (ensembliste et multiensembliste) coïncident.

Nous allons commencer par une analyse sémantique du contre-exemple de base (de la figure 6) à la confluence de PN_2 . L’interprétation de tous les réseaux de la figure 6 est l’union disjointe de quatre liens axiome. Soyons précis : l’ensemble qui interprète les deux réseaux obtenus en éliminant la coupure de départ (et donc aussi le réseau de départ) est

$$\{(x, 0), (x, 0), (y, 1), (y, 0), (z, 0), (z, 1), (t, 1), (t, 1) : x, y, z, t \in |\mathcal{A}|\}.$$

Il apparaît clairement que ce que la sémantique “voit”, c’est la position des liens axiomes dans les boîtes. En l’occurrence, quatre positions sont possibles : 00, 10, 01, 11. L’interprétation sémantique rend compte des réseaux de la figure 6 comme d’une juxtaposition (ou plutôt “superposition”) de quatre liens axiome, chacun avec sa position. Quant à savoir si le $\&$ de la boîte B_1 a été effectué avant ou après le $\&$ de la boîte B_2 , la sémantique ne dit absolument rien à ce sujet. Ce qui suggère que ces deux opérations peuvent être effectuées “en même temps”.

Figure 6. Non confluence de PN_2 .

Notre idée fut alors celle de trouver un moyen pour exprimer ce “parallelisme” dans la syntaxe. Idée d’autant plus alléchante qu’elle laisse espérer que sa mise en oeuvre fournisse une unique forme normale pour le réseau avec coupures de la figure 6.

Nous introduisons la notion de “multiboîte” : il s’agit d’une boîte additive avec plusieurs portes $\&$, qui présente (à priori) de multiples intérêts. D’une part l’intuition suivant laquelle une boîte additive indique la superposition de plusieurs processus de calcul est préservée, et d’autre part on peut espérer représenter, à l’aide de cette notion, “l’atemporalité” des règles $\&$ (qui est flagrante dans la sémantique). La séquentialisation ne pose aucun problème (grâce à la présence des boîtes), et la modification des e.e.r. que ce changement impose paraît gérable. Et bien, allons-y...

5.1.1. DÉFINITION. (structure de preuve) Une structure de preuve G est définie comme dans 1.2.1 (pour les liens ax , $\wp$, $\otimes$, $\oplus$, $\&$, $coad$ de $MALL$), avec quelques différences :

- tout lien $\&$ a 2^n prémisses (pour un certain $n \geq 1$) et une conclusion. Parmi ces prémisses, il y en a 2^{n-1} de même type (notons-le C), les autres 2^{n-1} étant aussi de même type (notons-le D). La conclusion du lien $\&$ sera de type $C\&D$ si toutes les arêtes de type C (resp. D) sont les prémisses gauches (resp. droites) du lien.
- tout lien $coad$ a 2^n prémisses (pour un certain $n \geq 1$) et une conclusion, toutes de même type
- à tout lien m qui est un lien $\&$ avec 2^n prémisses, est associé un (unique) sous-graphe B de G , tel que une parmi les conclusions de B est la conclusion de m et toute autre conclusion de B (il pourrait n’y en avoir aucune) est conclusion d’un lien $coad$ ayant 2^n prémisses ou bien d’un autre lien $\&$ ayant 2^n prémisses. B est toujours appelé une boîte additive ; m est une des portes $\&$ de B
- à tout lien $coad$ m qui a 2^n prémisses est associée une boîte additive B , telle que une parmi les conclusions de B est conclusion de m ; on dit que m est une porte $coad$ de B .

La condition d’emboîtement est toujours la même : deux boîtes quelconques sont disjointes, ou bien l’une d’entre elles est contenue dans l’autre.

5.1.2. DÉFINITION. (réseau) Un réseau se définit comme dans 1.3.7, avec les modifications évidentes qui s’imposent. Nous noterons $MALL$ l’ensemble de ces réseaux.

5.1.3. REMARQUE. Il est bien clair que tout réseau suivant la définition 1.3.7, qui ne fait intervenir que des connecteurs multiplicatifs et additifs, est un réseau de *MALL*. Et il est tout aussi clair qu'il existe des réseaux de *MALL* qui ne sont pas des réseaux dans le sens de la définition 1.3.7. Par exemple, le dernier réseau de la figure 6. Remarquons que ce réseau semble à juste titre pouvoir aspirer à l'appellation de "unique forme normale" du réseau de départ.

5.1.4. REMARQUE. La séquentialisation (le théorème 1.3.10 pour *MALL*) se prouve dans le fragment multiplicatif et additif de LL_2 , dans lequel on a modifié la règle d'introduction du $\&$: on aura une règle avec 2^n sous-preuves prémisses ($n \geq 1$), que nous allons décrire. Soient $D_1, \dots, D_n$ (resp. $E_1, \dots, E_n$) des (occurrences de) formules. Il y a une bijection entre les entiers $\{1, \dots, 2^n\}$ et les multiensembles $C_1, \dots, C_n$, où $C_i \in \{D_i, E_i\}$ (et $i \in \{1, \dots, n\}$). Cette bijection va associer à tout entier l un multiensemble $C_1^l, \dots, C_n^l$.

Pour tout $l \in \{1, \dots, 2^n\}$, soit α_l une preuve dont le séquent conclusion est $\vdash \Gamma, C_1^l, \dots, C_n^l$. La règle $\&$ de notre calcul des séquents permet d'obtenir à partir de ces 2^n preuves une preuve α de conclusion $\vdash \Gamma, D_1 \& E_1, \dots, D_n \& E_n$.

La même remarque que ci-dessus s'applique : si on appelle LL_{ma} ce calcul des séquents, toute preuve (multiplicative et additive) de LL_2 est une preuve de LL_{ma} , et il existe des preuves de LL_{ma} qui ne sont pas des preuves de LL_2 . Par exemple, il y a une séquentialisation du dernier réseau de la figure 6 qui est une preuve de LL_{ma} sans être une preuve de LL_2 .

Nous appliquerons dans la suite les mêmes conventions que dans les chapitres précédents, sans systématiquement redéfinir des notions qui s'étendent de façon évidente à notre nouveau cadre. Notamment, la notion de composante d'une boîte se généralise de façon évidente à *MALL*, et la définition de liens, arêtes ou boîtes séparés par une boîte donnée est alors la même que celle de la définition 1.3.1.

5.2 La normalisation des réseaux de *MALL*

Il est certes tentant de définir la normalisation exactement comme au chapitre 2. C'est ce que nous ferons, avec les modifications qui s'imposent.

Mais l'exemple de la figure 6 donne aussi envie de définir une autre étape de normalisation : celle qui permet de passer des deux formes normales (de PN_2) au nouveau réseau (de *MALL*). Est-ce licite ? Une telle étape

préserve la sémantique dénotationnelle (par construction!), la propriété de séparation du chapitre 3 (proposition 5.3.4), et nous montrerons aussi qu'en la rajoutant aux étapes usuelles de normalisation, on obtient une réécriture qui est fortement normalisante (théorème 5.3.8).

5.2.1. DÉFINITION. (normalisation dans $MALL$). Les coupures sont celles de la définition 2.2.1, si ce sont des coupures de notre fragment. Nous allons définir cinq e.e.r. : $[ax]$, $[\wp / \otimes]$, $[\& / \oplus]$, $[ccad]$, $[r(\&)]$.

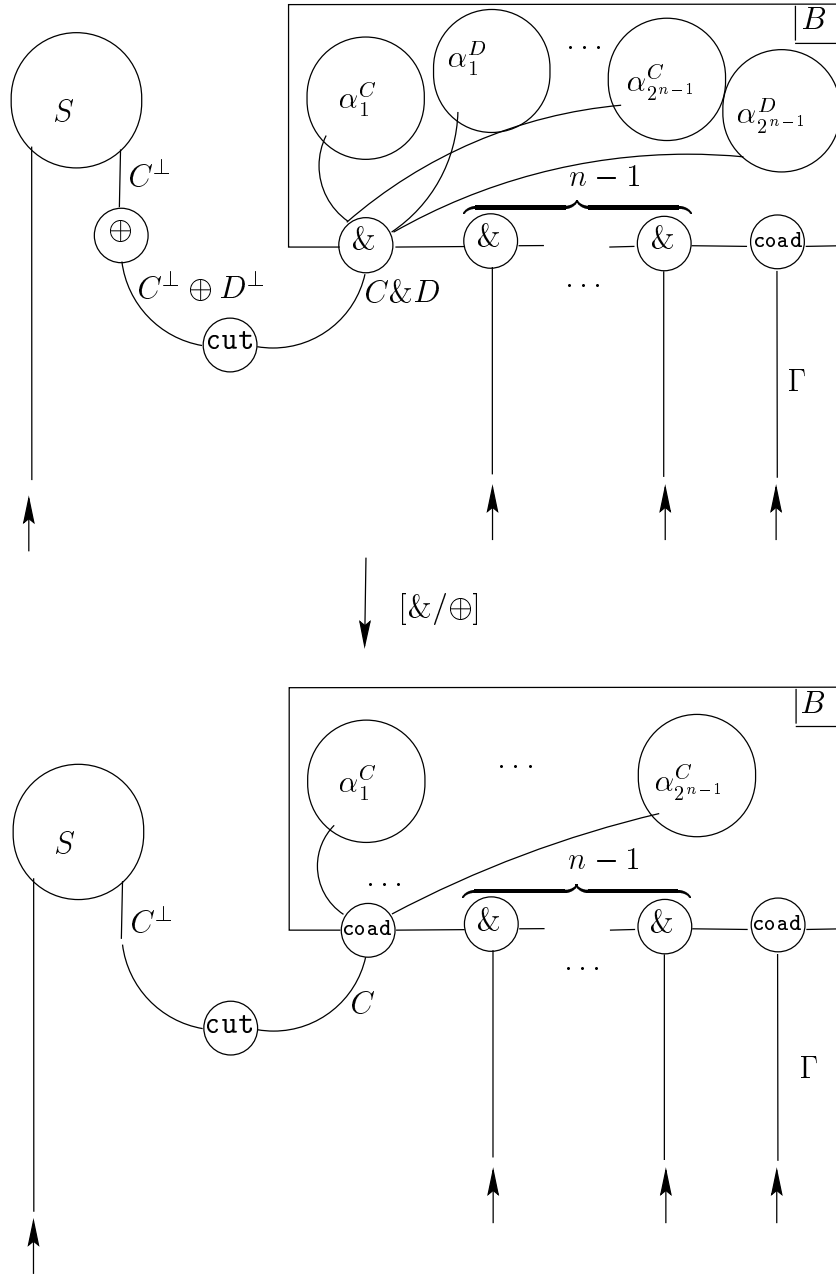
Les étapes $[ax]$ et $[\wp / \otimes]$ sont définies comme dans 2.2.3 (c.à.d. comme dans [Girard 87]).

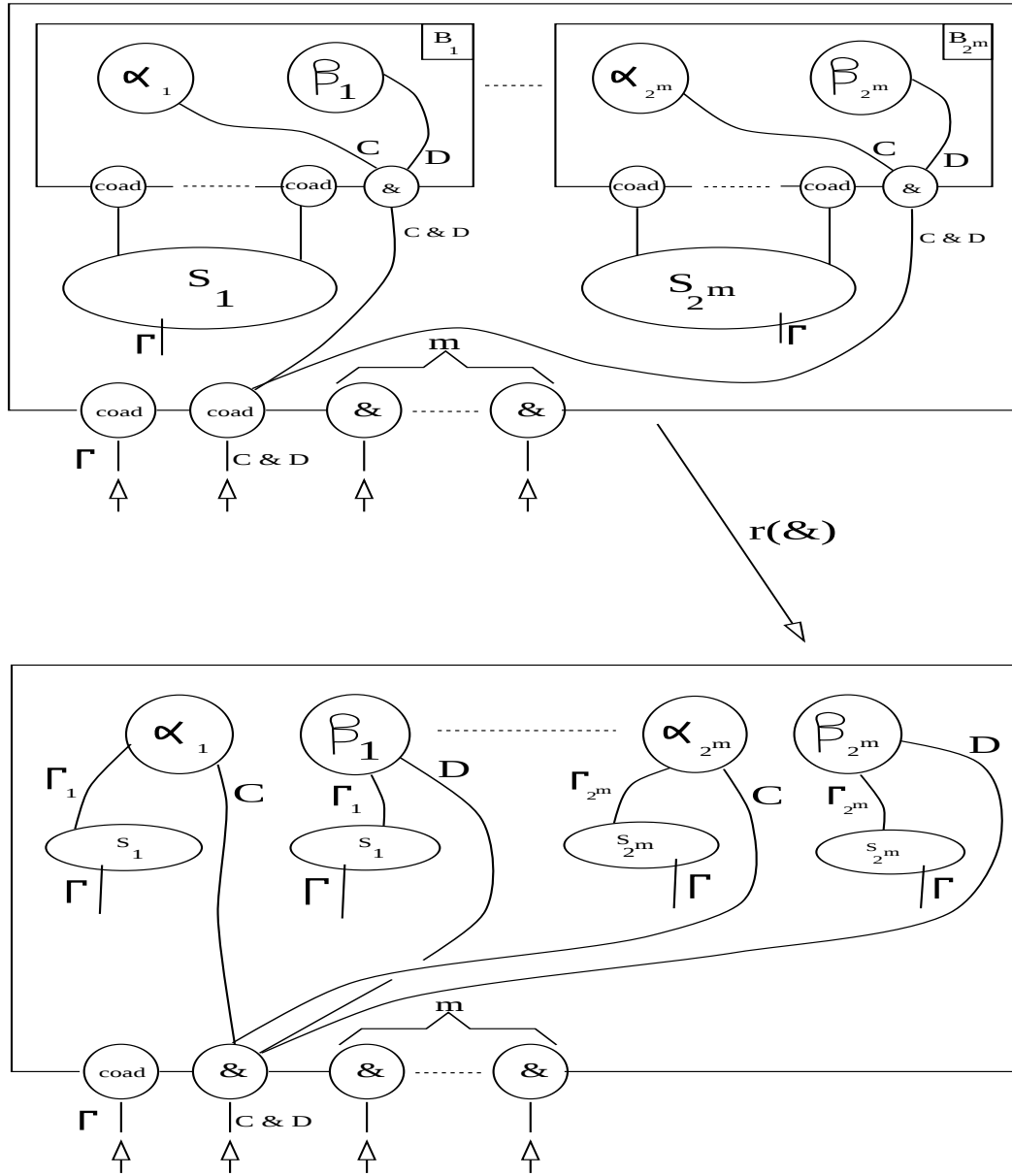
L'étape $[\& / \oplus]$ peut être “speciale” (le lien $\&$ n'a que deux prémisses), et dans ce cas c'est l'étape de normalisation usuelle (avec ouverture de la boîte), ou bien “non spéciale” (le lien $\&$ a 2^n prémisses, avec $n \geq 2$) comme dans la figure 7.

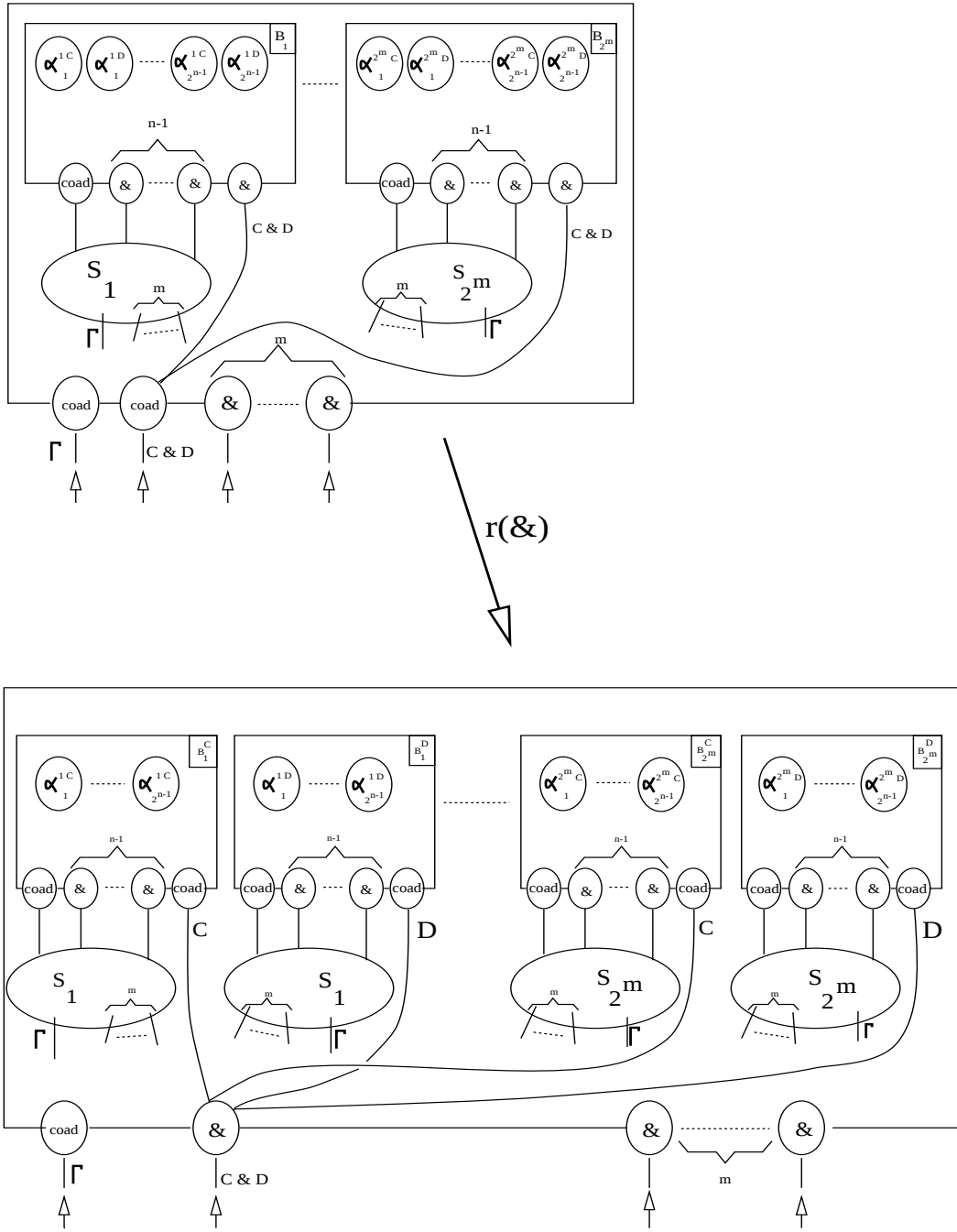
L'étape $[ccad]$ se définit comme dans 2.2.3, à la différence que le réseau S de la figure 4 du chapitre 2 va être copié 2^n fois, si la boîte additive B -attractive pour la coupure que l'on réduit contient 2^n composantes. (On reprend ici la terminologie des chapitres précédents : composantes, boîte attractive).

L'étape $[r(\&)]$ n'est pas liée à un lien coupure. Elle s'applique lorsque toutes les prémisses d'un lien n de type $coad$ d'une boîte sont conclusions d'un lien $\&$. Dans la version que nous considérons ici, nous supposons que les liens $\&$ dont la conclusion est prémisses de n sont portes de boîtes (différentes deux à deux par correction) qui contiennent toutes le même nombre de portes $\&$. Ce cas est plus simple à gérer (ne serait-ce que pour les notations!), mais rien ne s'oppose à une généralisation de cette étape.

L'étape peut être spéciale (comme dans la figure 8) ou bien non spéciale (comme dans la figure 9). Cette étape a (en général) comme effet celui de faire augmenter (considérablement...) le nombre de liens du réseau.

Figure 7. L'e.e.r. $[\&/\oplus]$ non spéciale ($n \geq 2$)

Figure 8. L'étape élémentaire de réduction $[r(\&)]$ spéciale.

Figure 9. L'étape élémentaire de réduction $[r(\&)]$ non spéciale.

5.2.2. REMARQUE. L'étape $[r(\&)]$ est en fait une façon de superposer des liens $\&$, qui a comme effet celui de rendre plus souple la relation entre une boîte et sa (ou ses) porte(s) $\&$.

A présent, une boîte n'est plus qu'une façon d'indiquer que l'on "superpose" un certain nombre de preuves.

Remarquons aussi que les deux formes normales (dans PN_2) du réseau de la figure 6 sont équivalentes "par renversement", c.à.d. qu'en appliquant un renversement à l'une d'entre elles on obtient l'autre. Il sera beaucoup question de cette opération de renversement dans le chapitre 15. Ces deux réseaux se réduisent, par une étape $[r(\&)]$ au réseau de la figure 6 contenant une multiboîte. La représentation de ces deux objets au moyen d'une multiboîte peut donc aussi être vue comme un représentant canonique dans la classe d'équivalence induite par l'opération de renversement.

5.2.3. REMARQUE. L'étape de réduction $[r(\&)]$ ne modifie pas les notions d'ancêtre et de résidu d'une boîte définies dans 2.2.6. Insistons sur le fait qu'il ne faut pas penser à une boîte comme à un sous-réseau mais plutôt comme à une arête (un peu spéciale, nous en convenons) du graphe. C'était d'ailleurs déjà l'optique de la définition 2.2.6, puisqu'une boîte et ses résidus pouvaient ne pas avoir le même contenu.

Pour tous les liens logique *qui ne sont pas* des liens $\&$, la définition 2.2.6 est toujours valable.

Ce n'est pas le cas des liens $\&$: ce sont les seuls liens logiques qui peuvent avoir *plusieurs* ancêtres (après une étape $[r(\&)]$).

5.3 Normalisation forte

Dans ce paragraphe, nous noterons $\rightarrow_u$ la réécriture "usuelle", c.à.d. que si $R \rightarrow_u R'$, alors il existe un lien coupure x de R tel que $R \xrightarrow{[x]} R'$. Nous noterons $\xrightarrow{r(\&)}$ l'e.e.r. $[r(\&)]$, et nous écrirons $R \rightarrow R'$ pour indiquer que $R \rightarrow_u R'$ ou bien que $R \xrightarrow{r(\&)} R'$. Comme d'habitude, on notera $\rightarrow$ la clôture réflexive et transitive de $\rightarrow$.

La définition suivante est une variante de la définition de [Girard 87] de "taille d'un réseau", qui tient compte de notre nouvelle syntaxe.

Dans [Girard 87], cette notion est utilisée pour donner une preuve combinatoire du "petit" théorème de normalisation forte pour PN_2 : si on n'élimine pas les coupures contraction, alors l'élimination des coupures est fortement

normalisante. C'est dans un but similaire que nous définissons notre variante de "taille d'un réseau".

5.3.1. DÉFINITION. (taille d'un réseau). Soit R un réseau et soient $B_1 \dots B_n$ les boîtes de profondeur zéro dans R . Supposons qu'il y ait k arêtes de profondeur zéro dans R qui ne sont conclusions d'aucune porte de boîte. La taille de R , notée $s(R)$, est définie par :

- (1) si $n = 0$, alors $s(R) = 2^k$
- (2) si R est une boîte B avec m portes $\&$ et si $\alpha_1 \dots \alpha_{2^m}$ sont ses composantes, alors $s(B) := (s(\alpha_1) + \dots + s(\alpha_{2^m})) \cdot 2^m$
- (3) si $n \geq 1$, alors $s(R) := s(B_1) \cdot \dots \cdot s(B_n) \cdot 2^k$.

5.3.2. REMARQUE. (i) Cette variante de la définition de [Girard 87], tient compte du fait que les portes $\&$ d'une boîte ne sont pas liées à la boîte elle-même : dans l'évaluation de la taille, elles sont traitées comme des formules usuelles, et la boîte ne sert qu'à superposer (donc à faire la somme des tailles de) ses composantes.

(ii) La conclusion d'un lien *coad* d'un réseau R n'est pas prise en compte dans le calcul de $s(R)$, seules ses prémisses le sont.

La taille $s(R)$ définie ci-dessus ne dépend pas de la complexité des formules, et la proposition qui suit devrait donc pouvoir s'étendre sans problème au système $MALL_2$ obtenu en rajoutant à $MALL$ les quantificateurs du second ordre.

5.3.3. PROPOSITION. Soient R et R' réseaux t.q. $R \rightarrow R'$.

- (i) Si $R \xrightarrow{[x]}_u R'$ et $[x] \neq [ccad]$, alors $s(R') < s(R)$.
- (ii) Si $R \xrightarrow{[ccad]} R'$, alors $s(R') = s(R)$.
- (iii) Si $R \xrightarrow{[r(\&)]} R'$, alors $s(R') = s(R)$.

Démonstration : Il s'agit de le vérifier, étape par étape. Nous allons décrire les "changements de taille" qui prouvent la proposition.

Si $[x] = [ax]$, alors une taille $r \cdot 2 \cdot 2$ est remplacée par une taille r .

Si $[x] = [\emptyset / \otimes]$, alors une taille $s \cdot 2^6$ est remplacée par une taille $s \cdot 2^4$.

Si $[x] = [\& / \oplus]$ est spéciale, alors une taille $(\alpha_C + \alpha_D) \cdot 2 \cdot 2 \cdot t$ est remplacée par une taille $\alpha_C \cdot t$. On a supposé ici que les types des prémisses de la coupure étaient $C \& D$ et $C^\perp \oplus D^\perp$, et on a noté la taille de la composante α_C dont une conclusion de type C est prémisses du lien $\&$ effacé par la coupure, comme la

composante elle-même. Dans la suite de la preuve, nous utiliserons toujours ce genre de conventions.

Si $[x] = [\&/\oplus]$ n'est pas spéciale, alors une taille

$$(\alpha_1^C + \alpha_1^D + \dots + \alpha_{2^n-1}^C + \alpha_{2^n-1}^D) \cdot 2^{n-1} \cdot 2 \cdot 2 \cdot t$$

est remplacée par une taille

$$(\alpha_1^C + \alpha_1^D + \dots + \alpha_{2^n-1}^C + \alpha_{2^n-1}^D) \cdot 2^{n-1} \cdot t$$

(voir aussi la figure 7).

Si $[x] = [ccad]$, alors une taille

$$s \cdot (\alpha_1 + \dots + \alpha_{2^n}) \cdot 2^n \cdot t$$

est remplacée par une taille

$$(s \cdot \alpha_1 + \dots + s \cdot \alpha_{2^n}) \cdot 2^n \cdot t$$

Dans ce cas la taille est donc bien la même avant et après réduction.

Si l'étape de réduction est une étape $[r(\&)]$ spéciale, alors une taille

$$((\alpha_1 + \beta_1) \cdot s_1 \cdot 2 + \dots + (\alpha_{2^m} + \beta_{2^m}) \cdot s_{2^m} \cdot 2) \cdot 2^m \cdot t$$

est remplacée par une taille

$$(\alpha_1 \cdot s_1 + \beta_1 \cdot s_1 + \dots + \alpha_{2^m} \cdot s_{2^m} + \beta_{2^m} \cdot s_{2^m}) \cdot 2^{m+1} \cdot t$$

(voir aussi la figure 8). Dans ce cas la taille est donc bien la même avant et après réduction.

Si l'étape de réduction est une étape $[r(\&)]$ non spéciale, alors une taille

$$(s_1 \cdot 2^{n-1} \cdot 2 \cdot (\alpha_1^{1C} + \alpha_1^{1D} + \dots + \alpha_{2^n-1}^{1C} + \alpha_{2^n-1}^{1D}) + \dots + s_{2^m} \cdot 2^{n-1} \cdot 2 \cdot (\alpha_1^{2^m C} + \alpha_1^{2^m D} + \dots + \alpha_{2^n-1}^{2^m C} + \alpha_{2^n-1}^{2^m D})) \cdot 2^m \cdot t$$

est remplacée par une taille

$$(s_1 \cdot 2^{n-1} \cdot (\alpha_1^{1C} + \dots + \alpha_{2^n-1}^{1C}) + s_1 \cdot 2^{n-1} \cdot (\alpha_1^{1D} + \dots + \alpha_{2^n-1}^{1D}) + \dots$$

$$+ \dots + s_{2^m} \cdot 2^{n-1} \cdot (\alpha_1^{2^m C} + \dots + \alpha_{2^n-1}^{2^m C}) + s_{2^m} \cdot 2^{n-1} \cdot (\alpha_1^{2^m D} + \dots + \alpha_{2^n-1}^{2^m D})) \cdot 2^{m+1} \cdot t$$

(voir aussi la figure 9). Dans ce cas la taille est donc bien la même avant et après réduction. $\square$

La proposition suivante étend le théorème de séparation (3.2.2) à *MALL*.

5.3.4. PROPOSITION. (multiséparation.) Soient T, T' deux réseaux, t.q. $T \xrightarrow{\sigma} T'$. Soit $\mathbb{B}$ une boîte de T , et $\mathbb{B}'$ un résidu de $\mathbb{B}$ dans T' . Pour tout résidu $\mathbb{B}''$ de $\mathbb{B}$ dans T' t.q. $\mathbb{B}'' \neq \mathbb{B}'$, il existe une boîte additive B de T' qui sépare $\mathbb{B}'$ et $\mathbb{B}''$.

Démonstration : Il s'agirait de reprendre la démonstration du théorème 3.2.2 et de l'étendre au cas des étapes $[r(\&)]$, ce que nous ne ferons pas. Nous allons nous contenter des deux remarques suivantes, qui expliquent les raisons pour lesquelles la preuve de 3.2.2 s'étend aux réseaux de *MALL* :

- (i) A chaque fois qu'on duplique une boîte, tous ses résidus sont dans des composantes différentes d'une autre boîte
- (ii) Si deux boîtes B et B' sont séparées, alors deux résidus quelconques $\vec{B}$ et $\vec{B}'$ de B et B' sont encore séparés.

Ces deux propriétés sont toujours vérifiées par la réduction $\xrightarrow{r(\&)}$.

Une preuve plus détaillée s'obtient par induction sur σ . □

5.3.5. REMARQUE. La propriété de séparation va nous permettre de borner à priori la profondeur maximale des réduits d'un réseau R donné. En effet, puisque (voir aussi la remarque 3.2.3) tout résidu d'une boîte B de R ne peut pas être contenu dans un autre résidu de B (par la propriété de séparation), nous savons que la profondeur maximale de tout réduit de R ne pourra en aucun cas dépasser le nombre de boîtes de R .

Nous définissons ci-dessous une mesure sur un réseau R qui va nous permettre de prouver le théorème de normalisation forte (5.3.8). Cette mesure étant un n -uple, nous ferons référence à l'ordre lexicographique.

5.3.6. DÉFINITION. Soit R un réseau contenant n boîtes. Soit B une boîte de R . On définit :

$compl_{\&}(B) :=$ nombre de liens *coad* de B dont le type de la conclusion est une formule dont le connecteur principal est un $\&$.

$I := \{B : B \text{ est une boîte de profondeur } i \text{ dans } R\}$, et $compl_{\&}^i(R) := \sum_{B \in I} compl_{\&}(B)$.

$\|R\|_{\&} := (\text{nb de liens de profondeur } 0, compl_{\&}^0(R), \dots, \text{nb de liens de profondeur } n-1, compl_{\&}^{n-1}(R))$.

5.3.7. PROPOSITION. Soit R un réseau. On a :

- (i) Si $R \xrightarrow{[ccad]} R'$, alors $\|R'\|_{\&} < \|R\|_{\&}$

(ii) Si $R \xrightarrow{r(\&)} R'$, alors $\|R'\|_{\&} < \|R\|_{\&}$.

Démonstration : (i) Soit $i \geq 0$ la profondeur de la boîte B de R attractive pour l'étape $[ccad]$. Soient $B_i, \dots, B_1$ (avec B_k de profondeur $k-1$) les boîtes de R contenant B .

Le nombre de liens de R de profondeur j est constant, $\forall j < i$.

Par ailleurs (si elle existe) B_i , ainsi que toute boîte de R qui n'est pas contenue dans B_i , aura un unique résidu dans R' , et celui-ci aura les mêmes portes $coad$ et $\&$, avec des conclusions de même type. C'est vrai en particulier pour toute boîte de R de profondeur $j < i$, et donc $compl_{\&}^j(R') = compl_{\&}^j(R)$ $\forall j < i$.

Le nombre de liens de profondeur i dans R' est strictement inférieur au nombre de liens de profondeur i dans R , et donc $\|R'\|_{\&} < \|R\|_{\&}$.

(ii) Soit B la boîte de R dont le nombre de portes $\&$ augmente et soit $\vec{B}$ son résidu dans R' . Supposons que B soit de profondeur i ($i \leq n-1$) et contenue dans les boîtes $B_i, \dots, B_1$.

Le nombre de liens de R' de profondeur j , pour tout $j \leq i$, est le même que le nombre de liens de profondeur j de R .

Par ailleurs, de même que dans le cas (i), (si elle existe) la boîte B_i , ainsi que toute boîte de R qui n'est pas contenue dans B_i , aura un unique résidu dans R' , et celui-ci aura les mêmes portes $coad$ et $\&$, avec des conclusions de même type. C'est vrai en particulier pour toute boîte de R de profondeur $j < i$, et donc $compl_{\&}^j(R') = compl_{\&}^j(R)$ $\forall j < i$.

Par définition de $\xrightarrow{r(\&)}$, on a $compl_{\&}(\vec{B}) < compl_{\&}(B)$, et on peut en déduire $compl_{\&}^i(R') < compl_{\&}^i(R)$. D'où $\|R'\|_{\&} < \|R\|_{\&}$. $\square$

5.3.8. THÉORÈME. La règle de réécriture $\longrightarrow$ est fortement normalisante sur l'ensemble des réseaux de $MALL$.

Démonstration : On prouve que si $R \longrightarrow R'$, alors $(s(R'), \|R'\|_{\&}) < (s(R), \|R\|_{\&})$, suivant l'ordre lexicographique.

C'est une conséquence immédiate des propositions 5.3.3, 5.3.7, 5.3.4. $\square$

5.3.9. REMARQUE. Suite au théorème 5.3.8, on pourrait (naïvement) rêver de prouver la confluence locale de la réécriture $\longrightarrow$, et donc la confluence de $\longrightarrow$ (comme dans la preuve de la proposition 4.1.7).

Ceci est clairement impossible. Pour s'en convaincre, le lecteur peut considérer le réseau obtenu à partir d'une boîte additive B en coupant un réseau S sur la conclusion d'une porte $coad$ de B . Il y a en général plusieurs étapes $[ccad]$

associées à cette coupure : on peut dupliquer S tout entier ou seulement un sous-réseau de S . Il est bien clair que les deux choix possibles peuvent être irréconciliables.

Il n'est pas exclu que par des restrictions “ad hoc” sur l'e.e.r. $[ccad]$, on puisse parvenir à résoudre ce genre de conflit. Mais l'expérience montre qu'une telle attitude conduit fatalement à des solutions insatisfaisantes. On va donc plutôt tenter de tirer partie de la modification de la syntaxe proposée dans les paragraphes précédents (qui, elle, n'est pas ad hoc !) pour définir une procédure de normalisation confluente pour la réécriture $\longrightarrow$.

5.4 Structure des réseaux de *MALL*

Nous allons analyser la structure des réseaux de *MALL*. Plus précisément, nous associerons à tout réseau sa “décomposition en réseaux minimaux” (définitions 5.4.5, 5.4.10). C'est à l'aide de cette décomposition que nous définirons une procédure d'élimination des coupures confluente, dans le paragraphe suivant.

5.4.1 Rappels, conventions, notations

On va confondre, dans la suite, une arête d'un réseau avec son type. Dans *MALL* ceci est “suffisamment peu ambigu”, les seuls liens de *MALL* ayant des prémisses et des conclusions de même type étant les liens *coad*.

Nous utiliserons dans la suite la notion d'empire d'une formule, définie dans [Girard 87] et utilisée depuis dans plusieurs travaux (par exemple [Girard 95a]). Nous renvoyons à ces travaux pour une définition précise, mais on peut aussi tout simplement penser à l'empire de A dans le réseau R , noté $e(A)$, comme au plus grand sous-réseau de R ayant A parmi ses conclusions. Nous n'allons pas rappeler toutes les propriétés des empires, en supposant une certaine familiarité de la part du lecteur avec cette notion.

Nous utiliserons les expressions “lien terminal” et “dernier lien”. Un lien n est un dernier lien d'un réseau R lorsque la conclusion de n est une conclusion de R . Un lien n est un lien terminal de R lorsqu'il existe une séquentialisation de R dont la dernière règle est celle qui correspond au lien n . Il s'agit donc de deux notions *très* différentes.

REMARQUE. Tout lien $\wp$ ou $\oplus$ qui est un dernier lien d'un réseau R est aussi un lien terminal de R .

Soit R' un sous-réseau de conclusion Γ' de R . On notera $R \setminus R'$ le réseau avec hypothèses obtenu en substituant dans R le sous-réseau R' par un lien hypothèse ayant comme conclusions les arêtes Γ' .

Soit S un sous-graphe du réseau R . Si en effaçant dans R les liens et les arêtes de S on obtient un réseau, celui-ci sera noté $R \setminus S$.

5.4.2. REMARQUE. (i) Soit R un réseau contenant une coupure sur A . Alors $e(A) \cap e(A^\perp) = \emptyset$.

(ii) Soit R un réseau et R' un sous-réseau de R de profondeur 0 et de conclusions Γ . Si tous les chemins issus des formules de Γ ne traversent que des liens $\wp$ et $\oplus$, alors ces chemins traversent tous les liens de $R \setminus R'$ (qui sont donc tous des liens $\wp$ ou $\oplus$).

(iii) Soit R un réseau et R' un sous-réseau de R de conclusions Γ' . Si il existe un chemin issu de $A \in \Gamma'$ ne traversant aucune coupure et de dernière arrête D , alors il existe un sous-réseau de R ayant D parmi ses conclusions, et contenant toutes les formules de Γ' (en particulier $\Gamma' \subseteq e(D)$ et $R' \subseteq e(D)$).

5.4.3. DÉFINITION. Soit R un réseau, $n = cut$, $\otimes$ un lien de R de profondeur zéro dont les prémisses sont A_1 et A_2 .

On dira que le lien n est splittant pour (ou splitte) les coupures de R , lorsque tout lien coupure de R différent de n est un lien de $e(A_1) \cup e(A_2)$.

5.4.4. LEMME. (splitting des coupures) Soit R un réseau. Si il existe une coupure de profondeur 0 dans R , alors on a deux possibilités :

- (i) il existe un lien $\otimes$ splittant pour les coupures de R
- (ii) (i) n'est pas vérifiée et alors il existe un lien coupure splittant pour les coupures de R .

Démonstration : Soit $\sharp R :=$ nombre de liens contenus dans R . On procède par induction sur $\sharp R$, en utilisant le théorème de séquentialisation.

Si ax est un lien terminal de R , alors R est un axiome et il n'y a pas de coupures de profondeur 0 dans R .

Autrement, supposons que (i) ne soit pas vérifiée. Puisque aucun lien $\otimes$ n'est splittant pour les coupures de R , il ne peut pas y avoir de liens $\otimes$ parmi les liens terminaux de R .

- Si R est une boîte, alors il n'y a pas de liens coupure de profondeur 0.
- Si parmi les liens terminaux de R il y a un lien coupure c , alors R est obtenu à partir de R_1 et R_2 par coupure et c est splittant pour les coupures de R .

- Si parmi les liens terminaux de R il y a un lien $\oplus$ (resp. $\wp$), alors R est obtenu à partir du réseau R_1 en rajoutant un lien $\oplus$ (resp. $\wp$). Si il existe une coupure de profondeur 0 dans R alors elle est dans R_1 . Tout lien n qui est splittant pour les coupures de R_1 est splittant pour les coupures de R . Un tel lien existe par hypothèse d'induction, et c'est forcément un lien coupure puisque R ne satisfait pas (i). $\square$

5.4.5. DÉFINITION. (réseaux minimaux) On dira que le réseau R est minimal, lorsque il ne contient aucun lien $\otimes$ splittant pour les coupures de R .

5.4.6. LEMME. Soit R_1 (resp. R_2) un sous-réseau du réseau R de conclusion $\Gamma_1, C \otimes_1 D, A$ (resp. Γ_2, E), et supposons que R soit obtenu en connectant A et E par un lien $\otimes_2$ de conclusion $A \otimes_2 E$.

Si le lien $\otimes_1$ est terminal dans R , alors il est terminal dans R_1 .

Démonstration : Soient $e_1(C)$ et $e_1(D)$ (resp. $e(C)$ et $e(D)$) les empires de C et D dans R_1 (resp. dans R). On a $e_1(C) \subseteq e(C)$ et $e_1(D) \subseteq e(D)$. Le lien $\otimes_1$ étant un dernier lien de R_1 , la seule possibilité pour qu'il ne soit pas terminal dans R_1 c'est que $e_1(C)$ et $e_1(D)$ soient connectés par un lien $\wp$.

Puisque tout sous-graphe de R_1 qui est un sous-réseau de R est aussi un sous-réseau de R_1 , ce lien $\wp$ n'est ni un lien de $e(C)$ ni un lien de $e(D)$, ce qui voudrait dire que $\otimes_1$ n'est pas terminal dans R . $\square$

5.4.7. LEMME. (de la décomposition) Soit R un réseau de conclusions Γ , avec (au moins) n liens $\otimes$ terminaux, $n \geq 1$.

Il existe un unique ensemble de réseaux $E_R = \{R_1, \dots, R_{n+1}\}$ tel que R s'obtient en connectant les éléments de E_R par ces n liens $\otimes$.

Démonstration : Par induction sur n , en utilisant le lemme 5.4.6. $\square$

5.4.8. LEMME. Soit R un réseau dont tous les derniers liens sont des portes de boîte, des $\otimes$, des axiomes ou des coupures, et supposons qu'il existe une coupure de profondeur 0 dans R .

Si il existe un lien $\otimes$ splittant pour les coupures de R , alors il existe aussi un lien $\otimes$ terminal.

Démonstration : Soit $\otimes_0$ de conclusion $A_0 \otimes D_0$ le lien $\otimes$ splittant pour les coupures de R . Puisqu'il existe une coupure de profondeur 0, $\otimes_0$ est

à profondeur 0 dans R . Soit alors Φ_0 le chemin issu de $A_0 \otimes D_0$ et ayant comme dernière arête une conclusion de R . Puisque $\otimes_0$ est splittant pour les coupures de R , le chemin Φ_0 ne peut pas traverser de coupures. Le dernier lien traversé par Φ_0 est forcément un lien $\otimes$: notons-le $\otimes_1$ (il se pourrait que $\otimes_0 = \otimes_1$). Soit A_1 la dernière arrête traversée par Φ_0 avant $\otimes_1$, et $A_1 \otimes D_1$ la conclusion de ce dernier (si $\otimes_0 = \otimes_1$ on prendra $A_1 = A_0$ et $D_1 = D_0$). Si $A_1 \neq A_0$, alors $e(A_1)$ contient toutes les coupures de R .

Si $R = e(A_1) \cup e(D_1)$, alors $\otimes_1$ est le lien terminal cherché.

Autrement, il existe un lien $\wp$ entre une conclusion de $e(A_1)$ et une conclusion de $e(D_1)$. Comme avant, soit Φ_1 le chemin issu de ce lien $\wp$ et ayant comme dernier lien un lien $\otimes$ et dernière arrête une conclusion de R . Appelons $\otimes_2$ ce lien $\otimes$, A_2 la dernière arrête traversée par Φ_1 avant $\otimes_2$, et $A_2 \otimes D_2$ la conclusion de ce dernier. Dans tous les cas, $e(A_2)$ contient toutes les coupures de R , et $e(A_2) \supsetneq e(A_1)$. Si $R = e(A_2) \cup e(D_2)$, alors on a fini. Dans le cas contraire, on recommence.

Remarquons pour conclure que $\forall i \geq 1$ on a $e(A_{i+1}) \supsetneq e(A_i)$, ce qui implique que pour un certain $k \geq 1$ on aura forcément $R = e(A_k) \cup e(D_k)$. $\square$

5.4.9. COROLLAIRE. Si dans un réseau R il n'y a aucun lien $\otimes, \wp, \oplus$ terminal, alors R est minimal.

Démonstration : Si dans R il n'y a aucun lien $\wp, \oplus$ terminal, alors il n'y a pas non plus de dernier lien $\wp, \oplus$. Donc les derniers liens de R sont tous des axiomes, des portes de boîtes, des $\otimes$ ou des coupures.

Si il n'existe pas de coupure de profondeur zéro dans R , alors R est un axiome ou une boîte, et le corollaire est vrai.

Si il existe un lien coupure de profondeur zéro dans R , alors on peut appliquer le lemme précédent : si il existait un lien $\otimes$ splittant pour les coupures de R , alors il existerait aussi un lien $\otimes$ terminal. $\square$

Si parmi les derniers liens du réseau R il y a des liens $\wp$ ou/et des liens $\oplus$, alors on peut les effacer (ainsi que leurs conclusions), et ce jusqu'à l'obtention d'un sous-réseau de R dont tous les derniers liens sont des liens $\otimes, ax, cut$, ou bien des portes de boîte. Dans la définition qui suit, nous noterons $R^{\wp\oplus}$ un tel réseau.

5.4.10. DÉFINITION. (Décomposition de R en réseaux minimaux)

Soit R un réseau. On définit par induction sur le nombre de liens de R , l'ensemble $D(R)$ des réseaux minimaux associés à R :

- si R est un axiome, alors $D(R) := \emptyset$

- si il existe des liens $\oplus$ ou bien $\wp$ terminaux, alors $D(R) := D(R^{\wp \oplus})$
- si il n'y a pas de liens $\oplus$ ni $\wp$ terminaux, et si il existe exactement $n \geq 1$ liens $\otimes$ terminaux dans R , alors soit $\{R_1, \dots, R_n\}$ l'ensemble associé à R par le lemme de la décomposition (lemme 5.4.7). On pose $D(R) := D(R_1) \cup \dots \cup D(R_n)$
- si il n'y a aucun lien $\otimes, \wp, \oplus$ terminal, et il existe un lien coupure terminal, alors $D(R) := \{R\}$ (on sait que R est minimal par le corollaire 5.4.9)
- autrement, R est forcément une boîte, et $D(R) := \emptyset$.

5.5 Procédure de normalisation confluente

Grâce à la notion de réseau minimal, nous allons définir une procédure de normalisation confluente pour *MALL* (théorème 5.5.15).

Nous allons d'abord montrer que lorsque dans un réseau R il n'y a que des coupures *ccad* (et il y en a au moins une), il existe une boîte particulière que nous appellerons “extrémiste” (lemme 5.5.5).

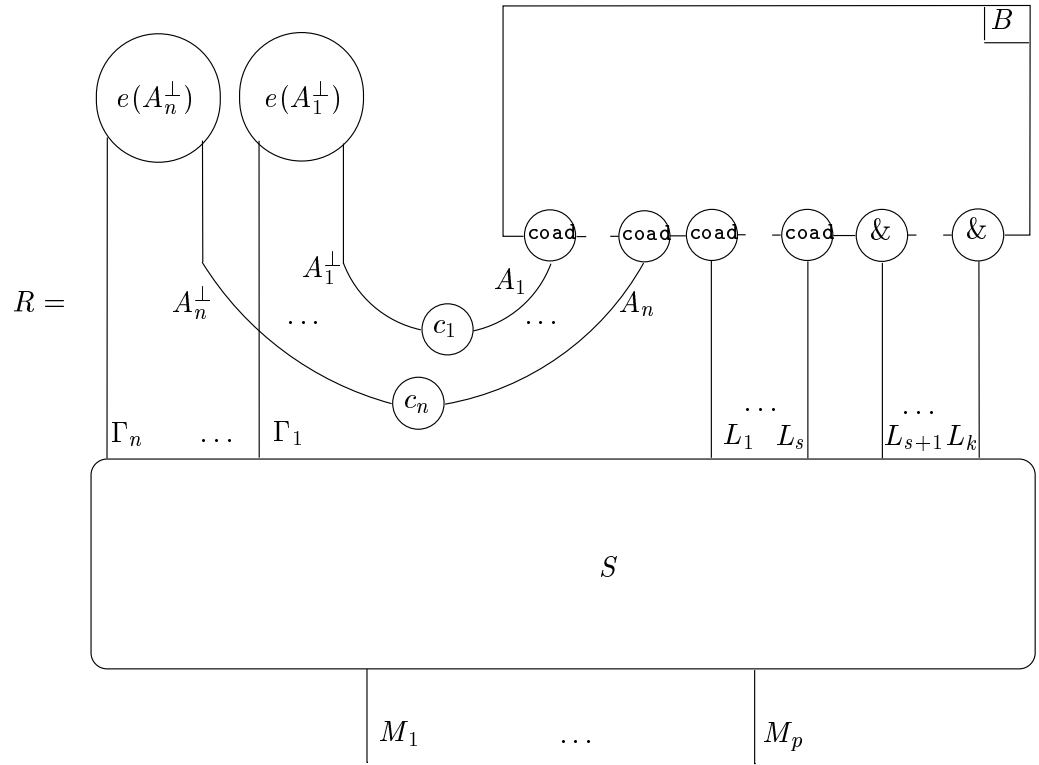
Ceci étant fait, nous n'aurons plus qu'à donner les définitions qu'il faut.

5.5.1. DÉFINITION. (forme canonique) Soit R un réseau minimal et B une boîte de profondeur 0 dans R telle que la conclusion d'une porte *coad* de B soit prémisses d'un lien coupure. On appellera forme canonique de R par rapport à B le réseau de la figure 10, où :

- les portes $\&$ de B ont comme conclusions $L_{s+1} \dots L_k$ et les portes *coad* ont comme conclusions $A_1 \dots A_n, L_1 \dots L_s$, avec $n \geq 1, k \geq 1, s \geq 0$
- $L_1, \dots, L_s$ ne sont pas prémisses d'un lien coupure
- pour tout $i \in \{1, \dots, n\}$ le graphe $e(A_i^\perp)$ est l'empire de $A_i^\perp$ dans R , dont les conclusions sont $A_i^\perp, \Gamma_i$
- S est un réseau avec hypothèses, le seul lien hypothèse de S ayant comme conclusions $L_1 \dots L_k, \Gamma_1 \dots \Gamma_n$, tandis que les conclusions de S sont $M_1 \dots M_p$ ($p \geq 1$), qui sont aussi les conclusions de R . On notera S à la fois le réseau avec hypothèses et le sous-graphe du réseau R .

Il nous arrivera de noter $R \setminus S$ le réseau obtenu à partir de R en effaçant le sous-graphe S de R .

5.5.2. DÉFINITION. Nous dirons d'un chemin issu d'un lien n d'un réseau R qu'il est **négligeable**, lorsqu'il ne traverse que des liens $\wp$ ou $\oplus$ (sauf éventuellement son lien initial et son lien terminal). Nous dirons qu'un chemin négligeable dont tous les liens sont à profondeur zéro est maximal, lorsque son lien terminal est un lien $\otimes$ ou bien *cut*.

Figure 10. Forme canonique de R par rapport à B

Soit R un réseau minimal contenant un lien $\otimes$ de profondeur 0 et de conclusion $A \otimes D$. Nous avons noté (dans la définition suivante) R_{AD} le sous-réseau de R obtenu en connectant les sous-réseaux $e(A)$ et $e(D)$ de conclusions respectives A, Γ_A et D, Γ_D , par un lien $\otimes$ de conclusion $A \otimes D$.

5.5.3. LEMME. (des $\otimes$ non splittants) Le sous-réseau R_{AD} de R satisfait une des deux conditions suivantes :

- (i) il existe un chemin négligeable maximal Φ issu de $A \otimes D$
- (ii) il existe $G \in \Gamma_A$ et $H \in \Gamma_D$ t.q. G et H sont prémisses d'un lien $\wp$, et il existe un chemin négligeable et maximal Φ issu de ce lien $\wp$.

Dans les deux cas, si Z est la dernière arrête de Φ , le sous-réseau $e(Z)$ de R contient l'arête $A \otimes D$ (et donc aussi R_{AD}).

Démonstration : Remarquons tout d'abord que toutes les conclusions de R_{AD} sont à profondeur 0, et que donc tout chemin issu de ces conclusions ne peut traverser aucun lien *coad* ou $\&$ sans avant avoir traversé une coupure. Remarquons ensuite que (par définition d'empire) toute formule de $\Gamma_A \cup \Gamma_D$ est active dans un lien $\wp$, ou bien c'est une conclusion de R .

Supposons que (i) ne soit pas satisfaite et prouvons (ii). Tout chemin issu de $A \otimes D$ est donc négligeable et non maximal. Soit alors Ψ le plus long chemin issu de $A \otimes D$ (dont la dernière arête sera forcément une conclusion de R).

On a deux possibilités :

cas 1 : Il n'y a, dans R , aucun lien $\wp$ dont les prémisses sont toutes les deux dans $\Gamma_A \cup \Gamma_D$.

Si Ψ ne traverse aucun lien $\wp$, alors par connexité des graphes de correction de R les formules de $\Gamma_A \cup \Gamma_D$ sont toutes conclusions de R , et le lien $\otimes$ de départ serait splittant pour les coupures de R (contredisant ainsi la minimalité de R).

Si Ψ traverse un lien $\wp$, alors par connexité des graphes de correction de R la prémisse des liens $\wp$ traversés par Ψ qui n'est pas une arête de Ψ est forcément parmi $\Gamma_A \cup \Gamma_D$. Et réciproquement, la conclusion de tout lien $\wp$ dont une prémisse est une formule de $\Gamma_A \cup \Gamma_D$ est une arête de Ψ . Ceci suffit pour pouvoir affirmer que $R \setminus R_{AD}$ ne contient que des liens $\wp$ et $\oplus$. Le lien $\otimes$ de départ serait alors splittant pour les coupures de R , contredisant ainsi la minimalité de R .

cas 2 : Supposons maintenant qu'il existe n liens $\wp$ ($n \geq 1$) entre $G_i \in \Gamma_A$ et $H_i \in \Gamma_D$ ($i \in \{1, \dots, n\}$). Il existe alors un sous-réseau R' de R de conclusions $A \otimes D, G_1 \wp H_1, \dots, G_n \wp H_n, \Gamma_A \cup \Gamma_D \setminus \{G_1, H_1, \dots, G_n, H_n\}$. Supposons par l'absurde que les chemins issus de $G_1 \wp H_1, \dots, G_n \wp H_n$ soient tous négligeables. Un raisonnement analogue à celui exposé ci-

dessus (faisant notamment toujours appel à la connexité des graphes de correction) montre que tout chemin issu d'une conclusion de R_{AD} est négligeable. Par (ii) de la remarque 5.4.2, $R \setminus R_{AD}$ ne contient alors que des liens $\wp$ ou $\oplus$. Il en suit que le lien $\otimes$ de conclusion $A \otimes D$ est splittant pour les coupures de R , contredisant ainsi la minimalité de R .

Remarquer que si (i) n'est pas vérifiée, alors le chemin Φ satisfaisant (ii) est de toutes façons issu d'une conclusion de R_{AD} , donc même dans ce cas $A \otimes D \in e(Z)$. $\square$

Soit R un réseau minimal ne contenant que des coupures *ccad*, et soit B une boîte de profondeur 0 dans R telle que l'une des conclusions des portes *coad* de B est prémisses d'un lien coupure. Considérons R dans sa représentation canonique par rapport à B .

5.5.4. LEMME. Si il existe une coupure dans S , alors il existe une coupure c dans S et une prémisses Q de c qui n'est pas conclusion d'un lien *coad*, et telle que $e(Q) \supseteq R \setminus S$.

Démonstration : Nous utiliserons ici les notations de la figure 10 et de la définition 5.5.1.

Remarquons tout d'abord que (par (ii) de la remarque 5.4.2), puisque S contient une coupure, il existe forcément un chemin négligeable maximal Φ_L issu d'une conclusion L de $R \setminus S$. On a alors deux possibilités pour S :

(i) il existe un chemin négligeable maximal Φ_L issu de la conclusion L de $R \setminus S$ dont le lien terminal est un lien coupure

(ii) il existe un chemin négligeable maximal Φ_L issu de la conclusion L de $R \setminus S$ dont le lien terminal est un lien $\otimes$.

cas (i) : Soit c le lien coupure qui est le lien terminal de Φ_L . Soit G la prémisses de c contenue dans Φ_L . Si $G = L$, alors (par définition d'empire) $G \notin \Gamma_i$, pour tout $i \in \{1, \dots, n\}$. On aura alors $G \in \{L_{s+1} \dots L_k\}$, c.à.d. que G est la conclusion d'une porte $\&$ de B : on a bien $e(G) \supseteq R \setminus S$ et G n'est pas conclusion d'une porte *coad*.

Si $G \neq L$, alors par définition de Φ_L , G est conclusion d'un lien $\wp$ ou $\oplus$ et (en appliquant (iii) de la remarque 5.4.2) $e(G) \supseteq R \setminus S$.

cas (ii) : Soit $A_0 \otimes D_0$ le lien terminal de Φ_L . On aura $e(C_0) \supseteq R \setminus S$, avec $C_0 \in \{A_0, D_0\}$. En appliquant le lemme des $\otimes$ non splittant (le lemme 5.5.3), on sait produire un chemin négligeable maximal Φ_0 dont la dernière arête C_1 sera prémisses d'un lien coupure ou bien d'un lien $\otimes_1$ de conclusion $A_1 \otimes D_1$. On aura $e(C_1) \supsetneq e(C_0)$.

Plus généralement, le lemme 5.5.3 permet de produire une suite de formules $C_0 \dots C_n$ (avec $C_i \in \{A_i, D_i\}$), et une suite de chemins négligeables maximaux $\Phi_0 \dots \Phi_n$ tels que :

- $e(C_n) \supsetneq \dots \supsetneq e(C_0) \supseteq R \setminus S$
- pour tout $i \in \{1, \dots, n\}$ la formule C_i n'est pas conclusion d'un lien *coad*
- Φ_i a comme dernière arête C_{i+1} pour $i \in \{1 \dots n\}$.

Les empires des C_i grandissant à chaque étape de notre construction, il faudra que pour n suffisamment grand la dernière arête C_{n+1} de Φ_n soit prémisses d'un lien coupure. L'arête C_{n+1} n'est pas conclusion d'une porte *coad*, et $e(C_{n+1}) \supseteq R \setminus S$. $\square$

5.5.5. LEMME. (des boîtes extrémistes) Soit R un réseau minimal avec, à profondeur 0, seulement des coupures *ccad* et au moins une telle coupure. Il existe une boîte B de R (de profondeur 0) telle que R puisse être représenté de façon canonique par rapport à B , où le réseau avec hypothèses S ne contient aucune coupure.

On dira qu'une telle boîte B est extrémiste.

Démonstration : On reprend, ici encore, les notations de la définition 5.5.1 et de la figure 10.

Pour toute boîte B de R de profondeur 0 ayant parmi les conclusions de ses portes *coad* une formule active dans une coupure (donc par rapport à laquelle on peut représenter canoniquement R suivant la définition 5.5.1), posons :

$$c_B := \{c_1, \dots, c_n\} \cup \{c' : c' \text{ est une coupure de } e(A_i^\perp), \text{ pour } i \in \{1, \dots, n\}\}.$$

Soit c_0 un lien coupure de profondeur 0 dans R . Appelons D_0 et $D_0^\perp$ les deux arêtes de R actives dans c_0 . Puisqu'il n'y a que des coupures *ccad* à profondeur 0, on peut supposer que D_0 est conclusion d'une porte *coad* de la boîte B_0 . On peut alors considérer la représentation canonique de R par rapport à B_0 . Si B_0 est extrémiste (c.à.d. qu'il n'y a pas de coupures dans S_0), alors on a fini.

Autrement, par le lemme 5.5.4, il y a dans S_0 un lien coupure c_1 dont une prémisses D_1 n'est pas conclusion d'une porte *coad*, et $e(D_1) \supseteq R \setminus S_0$.

La formule $D_1^\perp$ est donc forcément conclusion d'une porte *coad* de la boîte B_1 de S_0 . On peut donc représenter R canoniquement par rapport à B_1 , l'un des sous-réseaux coupés sur les conclusions des portes *coad* de B_1 étant $e(D_1)$. Puisque $e(D_1) \supseteq R \setminus S_0$, $\text{card}(c_{B_1}) > \text{card}(c_{B_0})$ ($c_1 \notin c_{B_0}$).

Comme dans la démonstration du lemme précédent, on peut ainsi construire une suite de boîtes de R qui sera forcément finie, et dont le dernier élément sera la boîte extrémiste cherchée. $\square$

5.5.6. COROLLAIRE. Dans le sous-graphe S de la représentation canonique du réseau R minimal par rapport à la boîte extrémiste B , il n'y a que des liens $\wp$ et $\oplus$.

Démonstration : Si il y avait un $\otimes$, alors ce $\otimes$ serait splittant pour les coupures de R . $\square$

Il s'agit maintenant de définir l'élimination des coupures de type *ccad*. La description d'une telle procédure est fatalement un peu lourde (essentiellement à cause des notations). Que le lecteur qui désire nous suivre jusqu'au théorème (5.5.15) s'arme donc d'un peu de patience.

Si les conclusions de toutes les portes $\&$ d'une multiboîte B sont conclusions du réseau R qui contient B , alors on peut définir le réseau obtenu à partir de R en substituant à B une quelconque de ses composantes. En généralisant cette substitution au cas de plusieurs boîtes, on aboutit à la notion de "tranche" d'un réseau, définie ci-dessous.

5.5.7. DÉFINITION. Soit R un réseau de conclusion $\Gamma, D_1^1 \& E_1^1, \dots, D_{k_1}^1 \& E_{k_1}^1, \dots, D_1^n \& E_1^n, \dots, D_{k_n}^n \& E_{k_n}^n$. Pour tout $j \in \{1, \dots, n\}$, soit B^j une boîte de R dont les conclusions des portes $\&$ sont $D_1^j \& E_1^j, \dots, D_{k_j}^j \& E_{k_j}^j$. Pour tout $j \in \{1, \dots, n\}$ et pour tout $l \in \{1, \dots, k_j\}$, nous associons l'entier $i_l^j = 1$ (resp. $i_l^j = 2$) aux 2^{k_j-1} composantes de B^j dont la conclusion est D_l^j (resp. E_l^j). On a donc une bijection entre les composantes de B^j et les suites $i_1^j \dots i_{k_j}^j$. Nous allons noter $B_{i_1^j \dots i_{k_j}^j}^j$ la composante de B^j de conclusions $C_1^j \dots C_{k_j}^j$ (avec $C_l^j \in \{D_l^j, E_l^j\}$), associée à la suite $i_1^j \dots i_{k_j}^j$. Soit $k = k_1 + \dots + k_n$. On appelle alors tranche de R par rapport à $B^1, \dots, B^n$, un quelconque des 2^k réseaux suivants :

$$R_{i_1^1 \dots i_{k_1}^1 \dots i_1^n \dots i_{k_n}^n} := R[B_{i_1^1 \dots i_{k_1}^1}^1 / B^1, \dots, B_{i_1^n \dots i_{k_n}^n}^n / B^n]$$

de conclusions $\Gamma, C_1^1, \dots, C_{k_1}^1, \dots, C_1^n, \dots, C_{k_n}^n$, avec $C_l^j \in \{D_l^j, E_l^j\}$.

5.5.8. DÉFINITION. Soit R un réseau minimal tel que toutes les coupures de profondeur zéro de R sont des coupures *ccad*.

On dira que la boîte extrémiste B de R est propre (dans R), lorsque dans la forme canonique de R par rapport à B , $S = \emptyset$ (en particulier toutes les portes $\&$ de B sont conclusions de R).

5.5.9. DÉFINITION. Soit R un réseau minimal n'ayant que des coupures *ccad* à profondeur 0.

On appelle sous-réseau propre de R le réseau $R^{\wp \oplus}$ obtenu à partir de R en effaçant tous les liens $\wp$ et $\oplus$ terminaux de R (ainsi que leurs conclusions), et ce jusqu'à l'obtention d'un sous-réseau de R dont tous les derniers liens sont des liens $\otimes$, *ax*, *cut*, ou bien des portes de boîte (voir aussi la définition 5.4.10).

5.5.10. REMARQUE. Dans $R^{\wp \oplus}$ toutes les boîtes extrémistes sont propres.

5.5.11. DÉFINITION. Soit R un réseau minimal t.q. toutes les coupures de profondeur 0 de R sont des coupures *ccad*, et soient $B^1, \dots, B^n$ les n boîtes extrémistes de R . Soient $D_1^j \& E_1^j, \dots, D_{k_j}^j \& E_{k_j}^j$ les $k_j \geq 1$ conclusions des portes $\&$ de B^j . Soit $R^{\wp \oplus}$ le sous-réseau propre de R . Les boîtes $B^1, \dots, B^n$ sont propres dans $R^{\wp \oplus}$.

On appelle réduit [*ccad*]-complet de R le réseau obtenu en substituant le sous-réseau $R^{\wp \oplus}$ de R par la multiboîte B ayant comme conclusions de ses portes $\&$ les formules $D_1^1 \& E_1^1, \dots, D_{k_1}^1 \& E_{k_1}^1, \dots, D_1^n \& E_1^n, \dots, D_{k_n}^n \& E_{k_n}^n$, et dont les composantes sont les tranches de R par rapport à ses n boîtes extrémistes. Plus précisément, les composantes de B sont les réseaux $R_{i_1^1 \dots i_{k_1}^1 \dots i_1^n \dots i_{k_n}^n}$ (définis dans 5.5.7) de conclusions $C_1^1, \dots, C_{k_1}^1, \dots, C_1^n, \dots, C_{k_n}^n$, où $C_l^j \in \{D_l^j, E_l^j\}$.

Soit B une boîte du réseau R , $F_1 \& G_1, \dots, F_l \& G_l$ les conclusions des portes $\&$ de B et A la conclusion d'une porte *coad* de B . Supposons que $F_1 \& G_1, \dots, F_l \& G_l$ soient des conclusions de R .

Soit R' un réseau de conclusions $A^\perp, \Gamma$, soient $B^1, \dots, B^n$ des boîtes de R' telles que pour tout $j \in \{1, \dots, n\}$ les conclusions des portes $\&$ de B^j sont $D_1^j \& E_1^j, \dots, D_{k_j}^j \& E_{k_j}^j$. Supposons que toutes ces formules soient conclusions de R' .

Appelons T le réseau obtenu en coupant B et R' sur A (remarquons au passage que les conclusions des portes $\&$ de $B, B^1, \dots, B^n$ sont toutes conclusions de T).

5.5.12. LEMME. Le réseau obtenu en réduisant la coupure *ccad* de T en dupliquant 2^l fois le réseau R d'abord, et en appliquant $k = k_1 + \dots + k_n$ e.e.r. $\xrightarrow{[r(\&)]}$ ensuite, est une boîte avec $k + l$ portes $\&$ de conclusions

$F_1 \& G_1, \dots, F_l \& G_l, D_1^1 \& E_1^1, \dots, D_{k_1}^1 \& E_{k_1}^1, \dots, D_1^n \& E_1^n, \dots, D_{k_n}^n \& E_{k_n}^n$, et qui a comme composantes les 2^{k+l} tranches de T par rapport à $B, B^1, \dots, B^n$.

Démonstration : C'est une conséquence des définitions. $\square$

5.5.13. PROPOSITION. Soit R un réseau minimal t.q. toutes les coupures de profondeur 0 de R sont des coupures *ccad*. Soient $B^1, \dots, B^n$ les n boîtes extrémistes de R , et $R^{\wp \oplus}$ le sous-réseau propre de R .

Le réduit $[ccad]$ -complet de R peut être obtenu à partir de la réduction $[ccad]$ usuelle et de $\xrightarrow{[r(\&)]}$ de la façon suivante :

- on fixe une boîte extrémiste B quelconque avec la forme canonique de R qui lui est associée (figure 10)
- on réduit (dans un ordre quelconque) les $[ccad]$ en dupliquant (dans $R^{\wp \oplus}$) les empires $e(A_1^\perp), \dots, e(A_n^\perp)$
- on applique les étapes $\xrightarrow{[r(\&)]}$ nécessaires aux portes *coad* de la boîte ainsi obtenue.

Démonstration : On applique à une boîte extrémiste quelconque le lemme précédent. $\square$

5.5.14. DÉFINITION. (la procédure de normalisation de MALL) Soit R un réseau. On va lui appliquer la procédure de normalisation suivante :

1. En appliquant la réécriture $\neg[ccad]$ à profondeur 0 (et sans utiliser d'étapes $[r(\&)]$), aboutir à un réseau R_0 ne contenant plus, à profondeur 0, que des coupures *ccad*
2. Décomposer R_0 en réseaux minimaux, suivant la définition 5.4.10
3. Associer à chaque réseau minimal de $D(R_0)$ son sous-réseau propre (suivant la définition 5.5.9)
4. Appliquer à chacun des réseaux obtenus de la sorte une étape $[ccad]$ -complète.

On aura à présent un réseau qui n'aura pas de coupures de profondeur 0. On pourra alors recommencer, en appliquant la procédure ci-dessus décrite à chacune des composantes des boîtes de profondeur 0.

5.5.15. THÉORÈME. (confluence) Soit R un réseau. La procédure de normalisation décrite ci-dessus n'utilise que des e.e.r. définies dans 5.2.1, et elle permet d'associer à R un unique réseau sans liens coupure.

Démonstration : Reprenons les notations de la définition. Le réseau R_0 est unique puisque la réécriture $\neg[ccad]$ est confluente (extension triviale de la

proposition 4.1.7 à *MALL*). L'unicité de la décomposition en réseaux minimaux est une conséquence de la définition 5.4.10 (elle-même conséquence du lemme 5.4.7). L'unicité du sous-réseau propre est une conséquence immédiate de la définition. L'unicité de toute étape *[ccad]*-complète découle aussi de la définition (5.5.11).

La procédure va forcément terminer, puisqu'on utilise seulement des e.e.r. définies dans 5.2.1 (par la proposition 5.5.13), et nous savons que celles-ci donnent lieu à une réécriture fortement normalisante (théorème 5.3.8). $\square$

5.6 Remarque sur le calcul *HC* de Girard

L'un des problèmes majeurs dans la recherche de la structure géométrique des preuves faisant intervenir le connecteur $\&$, est que celui-ci semble être indissociable d'une notion de superposition, qui reste vague et difficile à définir correctement. Nous avons montré que la superposition des liens $\&$ (et des liens $\&$ seulement) n'est pas complètement ingérable.

Une des idées essentielles du programme "ludique", lancé par J.-Y. Girard dans [Girard 98a], est la séparation des connecteurs logiques en deux groupes (positifs ou irréversibles d'une part et négatifs ou réversibles d'autre part), dont il sera d'ailleurs question dans les chapitres 14 et 15. Pour une première réalisation de son programme, Girard choisit le fragment multiplicatif et additif de LL, sans axiomes mais avec les constantes ([Girard 98b]).

L'aspect syntaxique de cette analyse utilise une notion qui s'apparente beaucoup à celle de multiboîte qui nous a occupés dans ce chapitre. Il n'est que rapidement question, dans [Girard 98b], des aspects syntaxiques du calcul "hyperséquentialisé" *HC*, mais il nous paraît clair que l'objet syntaxique qui découle de l'analyse sémantique de Girard est celui d'une "multiboîte étendue" à tous les connecteurs négatifs ($\&$ et $\wp$), et dont les portes *coad* ont comme conclusions toujours des formules positives. A vrai dire, à partir de la notion de multiboîte que nous venons de mettre au point, il n'est pas difficile de définir précisément une syntaxe de réseaux pour *HC*, qui n'est pas tout à fait celle de [Girard 98b] (il n'y aura pas de boîtes positives). Décrivons-la brièvement.

Il y a un lien axiome, représenté par une (multi)boîte dont la porte principale a comme conclusion la constante additive $\top$ (qui est une formule négative), et dont les portes *coad* ont toutes comme conclusion une formule positive.

Il y a un lien coupure ayant deux prémisses, une positive P et une négative $P^\perp$, et aucune conclusion.

A partir de 2^n réseaux dont toutes les conclusions sont positives (plus précisément le réseau R_i aura comme conclusion Γ, P_i , avec $i \in \{1, \dots, 2^n\}$) on peut construire un réseau avec une unique conclusion négative N , et un multienemble Γ de conclusions positives : il s'agit d'une multiboîte dont les composantes sont les 2^n réseaux de départ et la (seule) porte principale a comme conclusion N .

Il y a un lien positif ayant un nombre arbitraire $k \geq 1$ de prémisses, *toutes négatives et toutes conclusions de portes principales de multiboîtes* (notons-les $N_1, \dots, N_k$), et ayant comme conclusion la formule positive P égale à $N_1 \otimes \dots \otimes N_k$ “à quelques occurrences de $\oplus$ près”.

Tout réseau satisfait les deux propriétés suivantes :

1. il y a au plus un lien positif de profondeur 0
2. la conclusion de toute porte *coad* d'une multiboîte est une formule positive.

L'élimination des coupures pour un tel système se définit très simplement en utilisant les (généralisations appropriées des) e.e.r. $[\&/\oplus]$, $[\wp/\otimes]$, $[ccad]$. La propriété 2 permet d'exclure les “cas critiques” de coupures *ccad* (comme celui de la figure 6) : les prémisses d'un lien coupure ne peuvent en aucun cas être toutes les deux conclusions d'un lien *coad*.

La propriété 1 permet de définir une stratégie de normalisation purement séquentielle (appelée “paresseuse” dans [Girard 98b]). Remarquons que cette même propriété peut être vue comme une limite de la syntaxe que nous venons de définir, surtout si l'on pense à la souplesse de la syntaxe usuelle pour *MLL*.

Si nous avons décidé de ne pas développer la syntaxe décrite ci-dessus dans tous les détails, c'est parce que la théorie “ludique” est en évolution permanente, et on peut espérer aboutir à terme à une notion “plus souple” de réseau.

Chapitre 6

Sémantique des réseaux

Notre but étant celui d'établir des liens entre des objets "intentionnels" tels que les preuves de LL, et des objets "extensionnels" tels que les interprétations de ces mêmes preuves, nous nous intéresserons à des sémantiques *concrètes*. Il ne s'agit pas ici de savoir quels sont les axiomes abstraits que les structures adéquates à fournir une interprétation (ou une sémantique) doivent satisfaire, mais plutôt d'être en mesure de *calculer*, à partir d'un réseau de preuve, l'objet sémantique qui lui est associé. Pour ce faire nous utiliserons la notion d'expérience, introduite par J.-Y. Girard dans [Girard 87] pour la sémantique cohérente (ensembliste).

Nous ferons référence dans la suite aux sémantiques cohérentes ensembliste (cf. [Girard 87]) et multiensembliste (cf. [Girard 91b]), ainsi qu'à la sémantique relationnelle multiensembliste. Pour cette dernière, nous donnerons ci-dessous au lecteur les éléments nécessaires pour comprendre de quoi il s'agit, mais nous ne prouverons pas son invariance par élimination des coupures : le théorème 6.2.6, dont tout le monde est au courant mais dont personne n'a le courage (ou l'envie) d'écrire la preuve.

6.1 Préliminaires, conventions, notations

Dans nos investigations, nous nous bornerons au fragment propositionnel de LL, les quantificateurs faisant intervenir des complications techniques que nous choisissons d'ignorer (en tous cas dans cette thèse).

À vrai dire, après avoir donné la définition d'expérience pour LL propositionnelle (définition 6.2.1) nous nous placerons dans le fragment multiplicatif et exponentiel de LL (MELL), et ce non pas à cause de la sémantique des additifs (qui ne pose pas le moindre problème), mais de leur syntaxe : dans

la recherche de liens le plus étroits possibles entre syntaxe et sémantique, nous nous heurtons à l'absence d'une syntaxe vraiment satisfaisante pour les connecteurs additifs. C'est toujours motivés par cette même tentative de rapprochement, que nous avons tenté une amélioration de la syntaxe des additifs dans le cadre plus restreint de LL additive et multiplicative (MALL) dans le chapitre 5.

Tant que nous aurons affaire à tout le fragment propositionnel de LL (c.à.d. essentiellement dans la définition 6.2.1) nous ferons référence à la notion de réseau de la définition 1.3.7, à la seule différence que les liens contraction que nous considérerons ne seront plus forcément d'arité 2 (c.à.d. avec deux prémisses et une conclusion comme dans la définition 1.2.1) mais d'arité quelconque : tout lien $?co$ a désormais $k \geq 2$ prémisses et une conclusion. Ce changement induit quelques petites modifications dans les définitions de switching et de graphe de correction, que le lecteur n'aura aucun mal à apporter tout seul : elles s'imposent toutes, une fois établi que toutes les arêtes prémisses d'un lien $?co$ sont entre elles appariées. De même, il faudra modifier le calcul des séquents LL_2 et l'e.e.r. $[co]$ comme il se doit.

Une fois oubliés les additifs, la définition 1.3.7 fournit tout simplement un moyen de raccrocher chaque affaiblissement à un axiome ou bien à une boîte (exponentielle puisque ce sont les seules qui restent), de telle sorte que le réseau soit séquentialisable. Mais la fonction de saut n'interagit *jamais* avec le processus d'élimination des coupures (voir aussi la remarque 2.2.5). Le lecteur peut donc à sa guise se référer à la définition 1.3.7, ou bien oublier purement et simplement la fonction de saut et se référer à la définition 1.2.7, sans oublier cependant que tous les réseaux considérés sont séquentialisables dans le fragment multiplicatif et exponentiel de LL_2 . Nous noterons (de façon un peu abusive) $MELL$ l'ensemble de tous ces réseaux.

Nous focaliserons notre analyse sur les sémantiques multiensemblistes : dans le cas cohérent parce que les multiensembles semblent mieux s'adapter à nos objectifs (voir aussi le chapitre 8) et dans le cas relationnel tout simplement parce que la sémantique relationnelle ensembliste n'existe pas (elle ne fournit pas un invariant de l'élimination des coupures).

6.1.1 Conventions

Nous appellerons “type d'une arête a ” la formule de LL associée à l'arête a . Nous noterons entre accolades non pas (comme d'ordinaire) des ensembles mais des multiensembles (dans le sens de la définition 1.1.3) ; sauf mention explicite du contraire, nous allons toujours nous référer à des multiensembles. Pour la définition d'espace cohérent et pour l'interprétation en terme d'es-

paces cohérents des formules de LL, nous renvoyons à [Girard 87] et à [Girard 91b].

Voici venu le moment de dire quelques mots au sujet de la sémantique relationnelle. On pourrait la définir informellement comme “la sémantique cohérente sans cohérence”.

6.1.2. DÉFINITION. Soit $| \cdot |$ une fonction qui associe à toute variable propositionnelle un ensemble. Nous définissons l’extension de $| \cdot |$ à toutes les formules de LL de la façon suivante :

$$|\mathcal{A}^\perp| = |\mathcal{A}|$$

$$|\mathcal{A} \otimes \mathcal{B}| = |\mathcal{A}| \times |\mathcal{B}|$$

$$|\mathcal{A} \oplus \mathcal{B}| = |\mathcal{A}| \times \{0\} \cup |\mathcal{B}| \times \{1\}$$

$$|!\mathcal{A}| = M_f(|\mathcal{A}|)$$

où $M(|\mathcal{A}|)$ est le monoïde commutatif libre engendré par $|\mathcal{A}|$ (voir définition 1.1.3) et $M_f(|\mathcal{A}|)$ l’ensemble des multiensembles finis de $|\mathcal{A}|$.

Pour toute formule A de LL, $|\mathcal{A}|$ est la trame de l’espace interprétant la formule A dans la sémantique relationnelle. L’espace $\mathcal{A}$ est l’ensemble $\mathcal{P}(|\mathcal{A}|)$ des parties de $|\mathcal{A}|$.

Si A est une formule et E une occurrence de sous-formule de A , il nous arrivera de parler de la “complexité de $A \setminus E$ ” : il s’agit de $c_A - c_E$, où c_A (resp. c_E) est le nombre de connecteurs de A (resp. E).

Sauf mention explicite du contraire, tous les résultats et les définitions énoncés dans les chapitres qui suivent, seront valables à la fois pour la sémantique cohérente et pour la sémantique relationnelle (multiensemblistes toutes les deux).

6.1.3 Notations

On notera dans la suite :

- $a, b, c, \dots$ les arêtes d’un réseau,
- $A, B, C, \dots$ les types de ces arêtes,
- $\mathcal{A}, \mathcal{B}, \mathcal{C}, \dots$ les structures interprétant ces types,
- $|\mathcal{A}|, |\mathcal{B}|, |\mathcal{C}|, \dots$ les trames de ces structures,
- $\alpha, \beta, \alpha', \beta', \dots$ les arêtes de type atomique d’un réseau,
- $e, e', e_i, \epsilon, \epsilon_i, \dots$ les expériences,
- $\gamma, \gamma', \delta, \delta', \dots$ les résultats (ou les conclusions) de ces expériences
- $\Gamma, \Gamma', \Delta, \Delta', \dots$ des multiensembles de formules

- $\wp \Gamma$ la formule obtenue en effectuant le $\wp$ de toutes les occurrences de formules de Γ (dans un ordre donné). Par simplicité, nous noterons toujours $\wp \Gamma$ la structure qui interprète la formule $\wp \Gamma$.
- R_B le plus grand sous-réseau du réseau R qui est contenu dans la boîte exponentielle B de R ,
- $\text{card}(y)$ la cardinalité du multienemble y (ou de l'ensemble y si on a dit explicitement que y est un ensemble),
- $n[z]$ la répétition n fois de l'élément z : par exemple $\{n[z]\}$ est le multienemble contenant n occurrences de l'élément z ,
- $[R]_{\text{rel}}$ (resp. $[R]_{\text{coh}}$) l'interprétation du réseau R dans la sémantique relationnelle (resp. cohérente). Dans le cas de la sémantique cohérente, il nous arrivera d'écrire $[R]_{\text{coh}}$ pour préciser qu'il s'agit de l'interprétation de R dans la sémantique cohérente *multiensembliste*.

Nous utiliserons les notations usuelles pour la cohérence et l'incohérence (strictes ou larges). Lorsqu'il n'y aura pas d'ambiguïté, nous écrirons simplement $x \frown y$, au lieu de $x \frown y(\mathcal{A})$, pour $x, y \in |\mathcal{A}|$.

6.2 Les expériences

La définition d'expérience ci-dessous est l'extension à LL propositionnelle de la définition de [Girard 87], et ce pour les trois sémantiques que nous considérons dans cette thèse.

Dans l'énoncé suivant, un dernier lien d'une structure de preuve S , est un lien dont la ou les conclusion(s) est (sont) une (des) conclusion(s) de S . Un lien coupure de S est un lien conclusion si il est à profondeur (additive et exponentielle) zéro.

6.2.1. DÉFINITION. (expérience) Une expérience e d'une structure de preuve S est une application qui associe à toute arête a de S de type A un multienemble $e(a)$ d'éléments de $|\mathcal{A}|$, de telle sorte que si a est à profondeur (additive et exponentielle) 0, alors $e(a)$ est un singleton.

On définit une telle application e par induction sur S de la façon suivante :

- Si S est la structure de preuve vide, alors l'application vide est la seule expérience de S .

- Si S a parmi ses derniers liens un lien axiome n de conclusions les arêtes a_1 et a_2 de type A et $A^\perp$ respectivement, alors soit S' la structure de preuve obtenue à partir de S en éliminant le lien n et ses conclusions.

Si e' est une expérience de S' , alors toute application e définie par : $\forall a$ arête de S' $e(a) = e'(a)$ et $e(a_1) = e(a_2) = \{x\}$, avec $x \in |\mathcal{A}|$, est une expérience

de S .

- Si S a parmi ses derniers liens un lien $n = \wp$ (resp. $n = \otimes$) de conclusion c de type $C_1 \wp C_2$ (resp. $C_1 \otimes C_2$) et de prémisses c_1 de type C_1 et c_2 de type C_2 , alors appelons S' la structure de preuve obtenue à partir de S en éliminant le dernier lien n et sa conclusion.

Si e' est une expérience de S' , alors l'application e définie par : $\forall a$ arête de S' $e(a) = e'(a)$ et $e(c) = \{(x_1, x_2)\}$, avec $e'(c_1) = \{x_1\}$ et $e'(c_2) = \{x_2\}$ est une expérience de S .

- Si S a parmi ses derniers liens un lien $n = \oplus$ de conclusion c de type $C_1 \oplus C_2$ et de prémisses c_1 de type C_1 (resp. c_2 de type C_2), alors appelons S' la structure de preuve obtenue à partir de S en éliminant le dernier lien n et sa conclusion.

Si e' est une expérience de S' , alors l'application e définie par : $\forall a$ arête de S' $e(a) = e'(a)$ et $e(c) = \{(x_1, 0)\}$ (resp. $e(c) = \{(x_2, 1)\}$), avec $e'(c_1) = \{x_1\}$ (resp. $e'(c_2) = \{x_2\}$) est une expérience de S .

- Si S est une boîte additive dont la porte $\&$ a comme conclusion l'arête c de type $C_1 \& C_2$ et comme prémisses les arêtes c_1 de type C_1 et c_2 de type C_2 , et dont les portes coad ont comme conclusions les arêtes $a_1, \dots, a_n$ de types respectifs $A_1, \dots, A_n$ et comme prémisses les arêtes a_1^1 et a_1^2 de type $A_1, \dots, a_n^1$ et a_n^2 de type A_n , appelons S' la structure de preuve obtenue à partir de S en prenant le contenu de la boîte S . Si e_1 (resp. e_2) est une expérience de la sous-structure de preuve S_1 (resp. S_2) de S' ayant comme conclusions les arêtes $a_1^1, \dots, a_n^1, c_1$ (resp. $a_1^2, \dots, a_n^2, c_2$), alors :

(1) l'application définie par $\forall a$ arête de S' $e(a) = e_1(a)$ si a est une arête de S_1 et $e(a) = \emptyset$ si a est une arête de S_2 , et par $\forall i \in \{1, \dots, n\}$ $e(a_i) = e_1(a_i^1)$, $e(c) = \{(0, x_1)\}$ avec $e_1(c_1) = \{x_1\}$, est une expérience de S

(2) l'application définie par $\forall a$ arête de S' $e(a) = e_2(a)$ si a est une arête de S_2 et $e(a) = \emptyset$ si a est une arête de S_1 , et par $\forall i \in \{1, \dots, n\}$ $e(a_i) = e_2(a_i^2)$, $e(c) = \{(1, x_2)\}$ avec $e_2(c_2) = \{x_2\}$, est une expérience de S .

- Si S a parmi ses derniers liens un lien déréluction n de conclusion c de type $?C$ et de prémisses c' de type C , alors appelons S' la structure de preuve obtenue à partir de S en éliminant le dernier lien n et sa conclusion.

Si e' est une expérience de S' , alors l'application e définie par : $\forall a$ arête de S' $e(a) = e'(a)$ et par $e(c) = \{\{x\}\}$ où $e'(c') = \{x\}$, est une expérience de S .

- Si S a parmi ses derniers liens un lien contraction n d'arité k , de conclusion c de type $?C$ et de prémisses $c_1, \dots, c_k$ de type $?C$, alors appelons S' la structure de preuve obtenue à partir de S en éliminant le dernier lien n et sa conclusion.

Si e' est une expérience de S' , alors l'application e définie par : $\forall a$ arête de S' , $e(a) = e'(a)$ et par $e(c) = \{x_1 \cup \dots \cup x_k\}$, avec $e(c_i) = \{x_i\}$, est une expérience de S **si et seulement si** $x_1 \cup \dots \cup x_k \in |\mathcal{C}|$.

- Si S a parmi ses derniers liens un lien affaiblissement n de conclusion c , alors appelons S' la structure de preuve obtenue à partir de S en éliminant le dernier lien n et sa conclusion.

Si e' est une expérience de S' , alors l'application e définie par : $\forall a$ arête de S' , $e(a) = e'(a)$ et par $e(c) = \{\emptyset\}$ est une expérience de S .

- Si S a parmi ses derniers liens un lien coupure (qui sera donc à profondeur 0) n de prémisses c_1 et c_2 , alors appelons S' la structure de preuve obtenue à partir de S en éliminant le dernier lien n .

Si e' est une expérience de S' , alors l'application e définie par : $\forall a$ arête de S' , $e(a) = e'(a)$, est une expérience de S **si et seulement si** $e(c_1) = e(c_2)$.

- Si S est une boîte exponentielle dont la porte principale (resp. les portes auxiliaires) a comme conclusion l'arête c de type $!C$ et comme prémisses l'arête c' de type C (resp. ont comme conclusions les arêtes $a_1, \dots, a_n$ de types respectifs $?A_1, \dots, ?A_n$ et comme prémisses les arêtes $a'_1, \dots, a'_n$ de types respectifs $?A_1, \dots, ?A_n$), appelons S' la structure de preuve obtenue à partir de S en prenant le contenu de la boîte S .

Si $e_1, \dots, e_k$ sont k expériences de S' , alors l'application e définie par : $\forall a$ arête de S' $e(a) = e_1(a) \cup \dots \cup e_k(a)$, et par $\forall i \in \{1, \dots, n\}$ $e(a_i) = \{x_1^i \cup \dots \cup x_k^i\}$, où $\forall j \in \{1, \dots, k\}$ $e_j(a'_i) = \{x_j^i\}$, et $e(c) = \{\{y_1, \dots, y_k\}\}$, où $\forall j \in \{1, \dots, k\}$ $e_j(c') = \{y_j\}$, est une expérience de S **si et seulement si** $\forall i \in \{1, \dots, n\}$ $x_1^i \cup \dots \cup x_k^i \in |\mathcal{A}_i|$ et $\{y_1, \dots, y_k\} \in |\mathcal{C}|$.

Par ailleurs, l'application définie sur l'ensemble des arêtes de S par : $\forall a$ arête de S' $e(a) = \emptyset$ et $e(a_1) = \dots = e(a_n) = e(c) = \{\emptyset\}$ est une expérience de S .

- Si S est l'union disjointe des boîtes $S_1, \dots, S_n$ et si $\forall i \in \{1, \dots, n\}$ e_i est une expérience de S_i , alors l'application définie par : $\forall a$ arête de S_i $e(a) = e_i(a)$ est une expérience de S .

- Rien d'autre est une expérience.

Soit S une structure de preuve dont les arêtes conclusions sont $a_1, \dots, a_n$ de type $A_1, \dots, A_n$ respectivement, et soit e une expérience de R . Soit $\forall i \in \{1, \dots, n\}$ $e(a_i) = \{x_i\}$. On dira que $(x_1, \dots, x_n) \in |\mathcal{A}_1 \wp \dots \wp \mathcal{A}_n|$ est la conclusion (ou le résultat) de l'expérience e de R . On le notera aussi $x_1, \dots, x_n$.

6.2.2. REMARQUE. (i) Le lecteur attentif aura remarqué qu'il faudrait vérifier que la notion d'expérience est bien définie (lorsque la structure de

preuve S contient plusieurs derniers liens). Ce même lecteur aura aussi déjà compris que c'est bien le cas.

(ii) Soit e (resp. e') une expérience du réseau R (resp. R'), de résultat γ (resp. γ'). Si R et R' ont le même séquent conclusion (nous dirons tout simplement que R et R' ont même conclusion), alors l'égalité $\gamma = \gamma'$ induit entre les arêtes conclusions des deux réseaux une bijection évidente. Nous dirons de deux arêtes, l'une conclusion de R et l'autre conclusion de R' , que ce sont deux arêtes “correspondantes”, lorsque, justement, elles se correspondent via cette bijection.

(iii) Soit a une arête du réseau R et e une expérience de R . On parlera des éléments de $e(a)$ comme des étiquettes associées par l'expérience e à l'arête a .

(iv) Dans le cas de la sémantique relationnelle, il est important de remarquer que les conditions sur les étiquettes des conclusions des liens contraction et portes d'une boîte exponentielle sont *toujours* satisfaites. La seule condition non triviale dans le cadre relationnel concerne donc les liens coupures.

(v) Toute expérience cohérente (multiensembliste) d'un réseau R est une expérience relationnelle de R . (Ceci suppose bien sûr que si $\mathcal{X}_{coh}$ (resp. $\mathcal{X}_{rel}$) est l'espace qui interprète la variable propositionnelle X , alors $|\mathcal{X}|_{coh} = |\mathcal{X}|_{rel}$).

La définition suivante est celle de [Girard 87], que nous étendons au trois sémantiques que nous considérons.

6.2.3. DÉFINITION. Soit R un réseau de conclusion Γ . $[R] := \{\gamma \in |\wp \Gamma| : \text{il existe une expérience } e \text{ de } R \text{ de conclusion } \gamma\}$.

On dira que $[R]$ est l'interprétation de R dans la sémantique considérée. Il nous arrivera aussi de dire que $[R]$ est la sémantique de R .

6.2.4. REMARQUE. (i) Une conséquence immédiate de la définition ci-dessus est que pour tout réseau R , si on note $[R]_{coh}$ (resp. $[R]_{rel}$) la sémantique cohérente multiensembliste (resp. relationnelle) de R , on a $[R]_{coh} \subseteq [R]_{rel}$. Au sujet de la réciproque, voir la proposition 6.3.9.

(ii) Dans le cas particulier du fragment multiplicatif de LL, la réciproque du (v) de la remarque 6.2.2 est vraie : toute expérience relationnelle d'un réseau R est aussi une expérience cohérente de R . En particulier, dans ce cas (une fois fixée l'interprétation des atomes), l'interprétation de R est la même dans les trois sémantiques considérées dans cette thèse.

6.2.5. THÉORÈME. Si R est un réseau de conclusions Γ , alors $[R] \in \wp \Gamma$ (ici $\wp \Gamma$ est l'espace qui interprète la formule $\wp \Gamma$).

Démonstration : C'est évident dans le cas relationnel. Dans le cas cohérent, la preuve de [Girard 87] peut s'étendre sans difficultés au cas multiensembliste. $\square$

L'interprétation ci-dessus définie est invariante par élimination des coupures (c.à.d. qu'on a bien à faire à une sémantique *dénotationnelle* des réseaux de LL propositionnelle) :

6.2.6. THÉORÈME. Soit R un réseau et x une coupure de R . Si $R \xrightarrow{[x]} R'$, alors $[R] = [R']$.

Démonstration : Encore une fois, la démonstration de [Girard 87] s'étend sans difficultés au cas multiensembliste. La démonstration pour la sémantique relationnelle est omise. $\square$

6.2.7. REMARQUE. Dans [Girard 87], il y a une (fort jolie!) preuve de la propriété suivante pour le fragment multiplicatif de LL (et pour la sémantique cohérente) :

“une expérience d'un réseau donné est univoquement déterminée par sa conclusion”, autrement dit :

soit e (resp. e') une expérience de conclusion γ (resp. γ') du réseau R . Si $\gamma = \gamma'$, alors $e = e'$.

La remarque 6.2.4 (ii) montre que ce résultat s'étend clairement au cas de la sémantique relationnelle.

Dans sa démonstration, Girard utilise le critère de correction des longs voyages. E. Duquesne et J. Van de Wiele ont étendu ce résultat aux réseaux purs de LL (cf. [DuqVdW 94]), et il est vraisemblable que les réseaux de *MELL* satisfassent aussi la propriété en question. La démonstration de [DuqVdW 94] est assez laborieuse, et nous n'avons eu le courage de l'étendre à *MELL*.

Nous ne ferons d'ailleurs pas usage de cette propriété, qui n'est bien sûr pas satisfaite par la sémantique relationnelle (exercice pour le lecteur : trouver un contre-exemple!).

Convention : A partir de maintenant, nous ne considérerons plus que des réseaux de *MELL*.

Le processus d'élimination des coupures induit naturellement une relation d'équivalence (syntaxique) sur les réseaux. Par ailleurs, une sémantique dénotationnelle induit elle aussi une relation d'équivalence (sémantique) sur les réseaux. Le problème que nous nous posons est celui de comparer ces deux relations d'équivalence.

Une chose est sûre dès maintenant : si R est un réseau et $R \twoheadrightarrow R'$, alors R et R' sont équivalents, tant syntaxiquement que sémantiquement. On pourra donc restreindre, à juste titre, notre analyse aux réseaux sans coupures. Nous serons plus précis sur ce point au début du chapitre suivant.

6.3 Quelques propriétés des expériences

6.3.1. DÉFINITION. (restriction d'une expérience à un sous-réseau)

Soit R un réseau, e une expérience de R de conclusion γ et soit R_1 un sous-réseau de R . On va définir le multiensemble $e|_{R_1}$, dont les éléments sont des expériences de R_1 , par induction sur la profondeur p des conclusions de R_1 dans R .

Si $p = 0$, alors $e|_{R_1} = \{e_1\}$, où e_1 est l'expérience de R_1 définie par : pour toute arête a de R_1 $e_1(a) = e(a)$.

Soit $p + 1$ la profondeur des conclusions de R_1 dans R , soit B la boîte de R de profondeur 0 contenant R_1 et soit R_B le plus grand sous-réseau de R contenu dans B . Soit n la cardinalité de l'unique étiquette associée par e à l'arête conclusion de la pal de B et soient $e_1, \dots, e_n$ les n ($n \geq 0$) expériences de R_B à partir desquelles l'unique expérience du multiensemble $e|_B$ est obtenue (selon les modalités précisées dans la définition 6.2.1). On pose alors $e|_{R_1} = e_1|_{R_1} \cup \dots \cup e_n|_{R_1}$ (si $n = 0$, on aura $e|_{R_1} = \emptyset$).

Si $e|_{R_1} = \{e_1, \dots, e_l\}$, on notera $\gamma|_{R_1}$ le multiensemble $\{\gamma_1, \dots, \gamma_l\}$ où γ_i est la conclusion de l'expérience e_i de R_1 , $\forall i \in \{1, \dots, l\}$.

6.3.2. LEMME. Soit e une expérience de R , R_1 un sous-réseau de R et $e|_{R_1} = \{e_1, \dots, e_n\}$. Pour toute arête a de R_1 , $e(a) = e_1(a) \cup \dots \cup e_n(a)$.

Démonstration : Par induction sur la profondeur p dans R des arêtes conclusions de R_1 .

Si $p = 0$, alors $e|_{R_1} = \{e_1\}$ et par définition $e_1(a) = e(a)$ pour toute arête a de R_1 .

Soit $p + 1$ la profondeur dans R des arêtes conclusions de R_1 , et soit B la boîte de profondeur 0 dans R contenant R_1 : R_1 est sous-réseau de R_B . Si pour toute arête c conclusion de B , $e(c) = \{\emptyset\}$, alors $e|_{R_1} = \emptyset$, ce qui permet de conclure dans ce cas.

Autrement, soient $\epsilon_1, \dots, \epsilon_k$ ($k \geq 1$) les k expériences de R_B à partir desquelles l'expérience e est définie. Par définition d'expérience $e(a) = \epsilon_1(a) \cup \dots \cup \epsilon_k(a)$ pour toute arête a de R_1 , et par définition de restriction d'une expérience à un sous-réseau $e|_{R_1} = \epsilon_1|_{R_1} \cup \dots \cup \epsilon_k|_{R_1}$, c.à.d. que $\epsilon_1|_{R_1} \cup \dots \cup \epsilon_k|_{R_1} = \{e_1, \dots, e_n\}$. Posons alors $\forall i \in \{1, \dots, n\} \epsilon_i|_{R_1} = \{\epsilon_i^1, \dots, \epsilon_i^{l_i}\}$. En appliquant l'hypothèse d'induction au sous-réseau R_1 de R_B , on obtient $\forall i \in \{1, \dots, n\} \epsilon_i(a) = \epsilon_i^1(a) \cup \dots \cup \epsilon_i^{l_i}(a)$. Donc $e_1(a) \cup \dots \cup e_n(a) = \epsilon_1^1(a) \cup \dots \cup \epsilon_1^{l_1}(a) \cup \dots \cup \epsilon_k^1(a) \cup \dots \cup \epsilon_k^{l_k}(a) = \epsilon_1(a) \cup \dots \cup \epsilon_k(a) = e(a)$. $\square$

6.3.3. LEMME. Soit e une expérience du réseau R (non vide), et a une arête de R . On a $e(a) = \emptyset$ ssi il existe une boîte B de R contenant a t.q. si on appelle c la conclusion de la pal de B , alors on a $e(c) = \{n[\emptyset]\}$ pour un certain entier non nul n .

Démonstration : Soit B la plus petite boîte contenant a et t.q. si c est la conclusion de la pal de B , alors $e(c) \neq \emptyset$ (une telle boîte existe toujours car R et donc e n'est pas vide). Par définition d'expérience, le seul élément de $e(c)$ est $\emptyset$. $\square$

6.3.4. LEMME. Soit a une arête du réseau R de profondeur p et e une expérience de R . Soient $c_1, \dots, c_p$ de types $!C_1, \dots, !C_p$ les conclusions des p pal des boîtes de R contenant a .

Si $p \geq 1$ et $\forall i \in \{1, \dots, p\}$ la cardinalité de tous les éléments de $e(c_i)$ est égale à n_i (resp. si $p = 0$), alors $e(a)$ est un multiensemble de cardinalité $n_1 \cdot \dots \cdot n_p$ (resp. de cardinalité 1).

Démonstration : On procède par induction sur p .

Si $p = 0$, alors a est à profondeur 0 et par définition d'expérience $e(a)$ est un singleton.

Supposons que a soit à profondeur $p + 1$ dans R , et soient $c_0, c_1, \dots, c_p$ de types respectifs $!C_0, !C_1, \dots, !C_p$ les conclusions des $p + 1$ boîtes de R contenant a . Considérons la boîte B_0 de profondeur 0 dans R contenant a : soit c_0 (resp. c'_0) la conclusion (resp. la prémisse) de la pal de B_0 . Il existe alors n_0 expériences $e_1, \dots, e_{n_0}$ de R_{B_0} t.q. $\forall i \in \{1, \dots, n_0\}, e_i(c'_0) = \{x_i\}$. Puisque par le lemme 6.3.2 pour tout $j \in \{1, \dots, p\}$ on a $e(c_j) = e_1(c_j) \cup \dots \cup e_{n_0}(c_j)$, notre hypothèse permet d'affirmer que $\forall i \in \{1, \dots, n_0\}$ tous les éléments de $e_i(c_j)$ ont cardinalité n_j . On peut donc appliquer l'hypothèse d'induction à l'arête a de profondeur p du réseau R_{B_0} et à toute expérience e_i de R_{B_0} : $\forall i \in \{1, \dots, n_0\} e_i(a)$ est un multiensemble de cardinalité $n_1 \cdot \dots \cdot n_p$ si $p \geq 1$, et de cardinalité 1 si $p = 0$. De $e(a) = e_1(a) \cup \dots \cup e_{n_0}(a)$, on déduit

alors que $e(a)$ est un multiensemble de cardinalité $n_0 \cdot n_1 \cdot \dots \cdot n_p$ si $p \geq 1$ et de cardinalité n_0 si $p = 0$. $\square$

6.3.5. DÉFINITION. Soit a une arête de type A du réseau R et e une expérience de R . Soit $x \in e(a)$ (c.à.d. que $x \in |\mathcal{A}|$ est **une** des étiquettes associée à a par e). Pour toute occurrence de sous-formule C de A , on définit le multiensemble “projection multiensembliste de x sur C ”, que l’on note $|x|_C$, par induction sur la complexité de $A \setminus C$:

- si $C = A$, alors $|x|_C = \{x\}$
- si $E = C \otimes D$ ou bien $E = C \wp D$ (resp. $E = D \otimes C$ ou bien $E = D \wp C$) est une occurrence de sous-formule de A , alors $|x|_C = \{y \in |\mathcal{C}| : \exists z \in |\mathcal{D}| \text{ t.q. } (y, z) \in |x|_E\}$ (resp. $|x|_C = \{y \in |\mathcal{C}| : \exists z \in |\mathcal{D}| \text{ t.q. } (z, y) \in |x|_E\}$)
- si $D = ?C$ ou bien $D = !C$ avec D occurrence de sous-formule de A , alors $|x|_C = \{y \in |\mathcal{C}| : \exists z \in |\mathcal{D}| \text{ t.q. } y \in z \in |x|_D\}$.

6.3.6. REMARQUE. (i) Soit e (resp. e') une expérience du réseau R , soit a une arête de R de type A , C une occurrence de sous-formule de D et D une occurrence de sous-formule de A . Soit x (resp. x') une étiquette associée par e (resp. e') à l’arête a . Si $|x|_D = |x'|_D$, alors $|x|_C = |x'|_C$.

(ii) Soit a une arête de type A du réseau R et e une expérience de R . Soit $x \in |\mathcal{A}|$ **une** des étiquettes associée à a par e . On peut étendre la définition précédente au cas où C n’est pas une occurrence de sous-formule de A , en posant dans ce cas $|x|_C = \emptyset$.

(iii) Avec les notations de (ii), si $|x|_C \neq \emptyset$, alors C est une occurrence de sous-formule de A . Remarquons que la réciproque est fausse : on peut avoir $|x|_C = \emptyset$ avec C occurrence de sous-formule de A .

Dans les énoncés et les démonstrations du lemme et de la proposition suivants, on notera $|\mathcal{A}|_{coh}$ la trame de l’espace cohérent $\mathcal{A}$, et $[R]_{coh}$ (resp. $[R]_{rel}$) la sémantique cohérente multiensembliste (resp. relationnelle) du réseau R .

6.3.7. LEMME. Soit R un réseau de conclusion Γ et sans liens coupure, et soit e_r une expérience relationnelle de R de résultat γ .

Si $\gamma \in |\wp \Gamma|_{coh}$, alors il existe une expérience cohérente (multiensembliste) e_{coh} de R de même résultat que e_r .

Démonstration : Par induction sur une séquentialisation π_R de R .

Il s’agit en fait d’une simple vérification. On utilise la définition de $|\mathcal{A}|_{coh}$, et la propriété suivante des espaces cohérents : si $x' \subseteq x$ et $x \in \mathcal{A}$ (où $\mathcal{A}$ est un espace cohérent), alors $x' \in \mathcal{A}$.

Les deux cas plus compliqués sont ceux où la dernière règle de π_R est une contraction ou une règle bien sûr. Nous effectuons la vérification dans le cas où la dernière règle de π_R est une contraction, en laissant l'autre cas aux bons soins du lecteur.

Si la dernière règle de π_R est une contraction d'arité k , alors appelons c (resp. $c_1, \dots, c_k$) la conclusion (resp. les prémisses) du lien contraction n associé à la dernière règle de π_R . Soit R' le sous-réseau de R obtenu à partir de R en effaçant le lien n et sa conclusion c . Si $?C$ est le type de c et si $\Gamma = ?C, \Gamma'$, alors le multiensemble des types des conclusions de R' sera $?C, \dots, ?C, \Gamma'$. Soit x, γ' le résultat de e_R , où $e_R(c) = \{x\}$. La restriction $e_{R'}$ de e_R à R' est une expérience de R' de résultat $x_1, \dots, x_k, \gamma'$, où $e_{R'}(c_i) = \{x_i\}$ pour tout $i \in \{1, \dots, k\}$, et $x = x_1 \cup \dots \cup x_k$.

Puisque par hypothèse $\gamma \in |\wp \Gamma|_{coh}$, on aura en particulier $x = x_1 \cup \dots \cup x_k \in |\wp \Gamma|_{coh} = \mathcal{C}^\perp$, et donc $x_i \in |\wp \Gamma|_{coh}$ pour tout $i \in \{1, \dots, k\}$. On en déduit alors que $(x_1, \dots, x_k, \gamma') \in |\wp \Gamma|_{coh}$, ce qui permet d'appliquer à R' l'hypothèse d'induction : il existe une expérience cohérente e'_{coh} de R' de résultat $x_1, \dots, x_k, \gamma'$. Puisque $x \in |\wp \Gamma|_{coh}$, l'application e_{coh} des arêtes de R obtenue à partir de e'_{coh} en posant $e_{coh}(c) = \{x\}$ est une expérience cohérente multiensembliste de R de résultat γ . $\square$

6.3.8. REMARQUE. Une façon (informelle) d'exprimer la propriété que le lemme précédent met en évidence est la suivante : aucune des opérations définies dans 6.2.1 sur les étiquettes des arêtes d'un réseau ne peut “créer de la compatibilité”. Si (par hasard) pour une arête a de type A du réseau R et pour une expérience e_r relationnelle de R , $x \in e_r(a)$ est un élément de $|\mathcal{A}|_{coh}$, alors les étiquettes des arêtes de R qui sont “au-dessus de a ” sont, elles aussi, des éléments de la trame de l'espace cohérent qui interprète leur type.

L'énoncé de la proposition suivante nous a été suggéré par Thomas Ehrhard.

Fixons une interprétation relationnelle et une interprétation cohérente des variables propositionnelles du langage, telles que si $\mathcal{X}_{rel}$ (resp. $\mathcal{X}_{coh}$) est l'espace qui interprète la variable propositionnelle X , alors $|\mathcal{X}|_{coh} = |\mathcal{X}|_{rel}$.

6.3.9. PROPOSITION. Soit R un réseau de conclusion Γ . On a $[R]_{coh} = [R]_{rel} \cap |\wp \Gamma|_{coh}$.

Démonstration : Le théorème 6.2.6 permet de se restreindre au cas où R est un réseau sans coupures.

Le fait que $[R]_{coh} \subseteq [R]_{rel} \cap |\wp \Gamma|_{coh}$ est évident. Réciproquement, soit $\gamma \in [R]_{rel}$ le résultat de l'expérience relationnelle e_r de R . Si $\gamma \in |\wp \Gamma|_{coh}$, alors le lemme précédent permet d'affirmer qu'il existe une expérience cohérente (multiensembliste) e_{coh} de R de résultat γ : donc $\gamma \in [R]_{coh}$. $\square$

Remarquons qu'il découle de cette dernière proposition que si R et R' sont deux réseaux de mêmes conclusions, alors de $[R]_{rel} = [R']_{rel}$ on déduit $[R]_{coh} = [R']_{coh}$. Nous reviendrons sur les conséquences de cette proposition dans les remarques 7.1.4 et 10.2.5.

Chapitre 7

La question de l'injectivité

L'idée la plus naturelle lorsqu'on décide de comparer les relations d'équivalence (syntaxique et sémantique) dont il était question à la fin du paragraphe 6.2 du chapitre 6, c'est de chercher à montrer qu'elles coïncident. Nous verrons dans ce chapitre que c'est le cas pour le fragment multiplicatif de LL. Ceci n'étonnera pas grand monde, et la raison pour laquelle nous élevons au rang de théorème (le 7.3.4) ce qui n'est pas beaucoup plus qu'une remarque, c'est que les petits lemmes dont on a besoin pour le prouver nous indiquent la voie à suivre dans le cas (bien plus complexe) de *MELL*.

Si pour une certaine sémantique s et pour un certain ensemble E de réseaux, les relations d'équivalence syntaxique et sémantique coïncident, alors nous dirons que la sémantique s est **injective** pour l'ensemble E : deux réseaux sans coupures et différents auront deux interprétations différentes dans la sémantique s .

Nous noterons *MLL* l'ensemble des réseaux de *MELL* ne faisant intervenir que des liens axiome, coupure, par et tenseur. Puisque la différence entre deux quelconques parmi les trois sémantiques suivantes n'est détectable *que* en présence des exponentielles, l'injectivité pour *MLL* sera établie pour les trois : la sémantique cohérente ensembliste, la sémantique cohérente multiensembliste, et la sémantique relationnelle.

7.1 Position du problème

Pour bien poser le problème de l'injectivité, il faut faire quelques remarques préalables.

Tout d'abord il faut préciser ce qu'on entend par "interprétation" d'un

réseau (relativement à une sémantique donnée) : celle-ci est certes définie (essentiellement) dans [Girard 87], mais pas complètement, puisqu'elle dépend de l'interprétation des formules atomiques. Nous dirons que deux réseaux de mêmes conclusions sont sémantiquement équivalents lorsqu'ils ont même interprétation (dans la sémantique considérée), et ce *quelque soit* l'interprétation des atomes des types de leurs conclusions (voir définition 7.1.2).

Venons-en maintenant à l'équivalence syntaxique. Soit c un connecteur binaire de LL, et soit $c^\perp$ son dual. Les sémantiques que nous considérons identifient toutes l'axiome de conclusions $AcB, A^\perp c^\perp B^\perp$ et la preuve canonique de $AcB, A^\perp c^\perp B^\perp$ (à partir des axiomes de conclusions $A, A^\perp$ et $B, B^\perp$). On se réfère souvent à cette dernière preuve en disant que c'est la η -**expansion** de l'axiome (terminologie que nous héritons du λ -calcul). Nous n'avons donc aucun espoir de distinguer deux preuves qui ne diffèrent que par des η -expansions.

De la même façon, il est facile de voir que les sémantiques considérées ne pourront pas “voir” si la conclusion d'un lien affaiblissement ou contraction est prémisses d'un lien pax ou non. Et un affaiblissement ou une contraction dont la conclusion est prémisses d'une contraction est aussi “sémantiquement invisible”.

Pour bien poser le problème, il faut donc définir plus précisément les relations d'équivalence dont on veut montrer l'égalité.

7.1.1. DÉFINITION. Soit R un réseau quelconque de $MELL$. Nous dirons que R est **standard** lorsque :

- R est sans coupures
- toute conclusion d'un lien axiome de R est de type atomique
- si a est conclusion d'un lien $?w$ ou $?co$ de R , alors a n'est pas prémisses d'un lien pax ni d'un lien $?co$.

Il est bien évident qu'en effectuant les η -expansions nécessaires, en effaçant les liens qu'il faut, et en déplaçant (quand c'est possible) les liens $?w$ et $?co$ “hors des boîtes”, on peut associer à tout réseau sans coupures de $MELL$ un unique réseau standard.

Remarquons au passage que (η -expansions mises à part), un réseau standard n'est autre qu'un réseau de la “nouvelle syntaxe” définie par V. Danos et L. Regnier (par exemple dans [Regnier 92]).

Nous utilisons dans la définition suivante la confluence des réseaux de $MELL$ démontrée dans [Danos 90].

7.1.2. DÉFINITION. Soient R et R' deux réseaux de $MELL$ de même conclusion, soit s une sémantique de $MELL$, et notons $[T]_s$ l'élément de $\emptyset\Gamma$ associé par la sémantique s au réseau T de conclusion Γ .

Soit R_0 (resp. R'_0) l'unique réseau standard associé à l'unique forme normale de R (resp. R').

Nous dirons que R et R' sont **sémantiquement équivalents** lorsque $[R]_s = [R']_s$ (pour toute interprétation des atomes des types des conclusions de R). Nous écrirons simplement $[R]_s = [R']_s$ (ou $[R] = [R']$ lorsqu'il n'y aura pas d'ambiguïté).

Nous dirons que R et R' sont **syntactiquement équivalents** ou bien $\beta\eta$ -**équivalents** lorsque $R_0 = R'_0$. On écrira alors $R \simeq_{\beta\eta} R'$.

Nous sommes maintenant en mesure de formuler notre

7.1.3 Conjecture :

Soient R et R' deux réseaux de $MELL$ de même conclusion, soit s la sémantique cohérente ensembliste, ou bien la sémantique cohérente multiensembliste, ou encore la sémantique relationnelle de $MELL$.

$[R]_s = [R']_s$ ssi $R \simeq_{\beta\eta} R'$.

Autrement dit : la sémantique s est injective.

Remarquons qu'une des deux directions est conséquence immédiate des théorèmes déjà énoncés : si $R \simeq_{\beta\eta} R'$, alors $[R]_s = [R']_s$.

7.1.4. REMARQUE. La proposition 6.3.9 implique que lorsque $[R]_{rel} = [R']_{rel}$ on a aussi $[R]_{coh} = [R']_{coh}$. Il en découle que pour tout fragment F de LL, si la sémantique cohérente est injective pour F , alors la sémantique relationnelle l'est aussi.

Convention :

Dorénavant, tous les réseaux seront des réseaux standard de $MELL$. Ceci restera vrai dans tous les chapitres qui suivent, sauf (bien sûr) mention explicite du contraire.

La conjecture ci-dessus peut alors se reformuler de la façon suivante :

“Soient R et R' deux réseaux de même conclusion. Si R et R' sont sémantiquement équivalents, alors $R = R'$.”

7.2 Sous-graphes d'un réseau

Nous introduisons une terminologie dont on se servira dans la suite, et surtout dans les chapitres à venir.

7.2.1. DÉFINITION. À tout réseau R sont naturellement associés les graphes suivants :

- la structure de preuve ouverte, notée $SPO(R)$, obtenue à partir de R en effaçant tous les liens axiomes de R
- la structure de preuve linéaire (resp. la structure de preuve linéaire ouverte), notée $SPL(R)$ (resp. notée $SPLO(R)$ ¹), obtenue à partir de R (resp. de $SPO(R)$) en effaçant toutes les connexions entre les différentes portes d'une même boîte
- la structure de preuve partielle (resp. la structure de preuve partielle ouverte), notée $SPP(R)$ (resp. notée $SPPO(R)$) obtenue à partir de $SPL(R)$ (resp. $SPLO(R)$) en effaçant tous les liens pax (les portes auxiliaires des boîtes exponentielles) ainsi que leurs conclusions.

7.2.2. REMARQUE. (i) Soit a une arête du réseau R conclusion d'un lien porte auxiliaire ou dérélction m . Puisque la conclusion de tout lien contraction ou affaiblissement n'est pas active dans un lien porte auxiliaire (et puisque les conclusions de tout lien axiome de R sont atomiques), il existe un unique lien dérélction n de conclusion l'arête a' de R et tel que le chemin issu de a' et ayant comme arête terminale une des conclusions de R , traverse a après avoir traversé un certain nombre de liens porte auxiliaire. On dira que n est **le lien dérélction “au-dessus” de a** ou bien **“au-dessus” de m** . (On peut bien-sûr avoir $n = m$).

(ii) Soit a (resp. $a_1, \dots, a_k$) la conclusion (resp. les prémisses) d'un lien contraction m d'arité k . La position des règles structurelles dans un réseau standard permet d'affirmer que les arêtes $a_1, \dots, a_k$ sont conclusions de liens dérélction ou pax. (i) permet alors d'associer au lien m k liens dérélction $n_1, \dots, n_k$ (où n_i est la dérélction au-dessus de a_i) que l'on appellera **les k liens dérélction au-dessus de m** .

La définition de graphe d'une arête dans un réseau qui suit peut être donnée pour tous les réseaux de LL propositionnelle (pas forcément standard).

¹à ne pas confondre avec S.P.Q.R., s'il vous plaît

7.2.3. DÉFINITION. Soit R un pré-réseau et a une arête de R . Soit k_a la cardinalité de l'ensemble $\{n : n \text{ lien de } R \text{ t.q. il existe } \phi_a^n \text{ chemin direct de } R \text{ issu d'un lien axiome ou affaiblissement contenant } n \text{ et de conclusion } a\}$. Le graphe de a dans R , noté G_a^R ou G_a si il n'y a pas d'ambiguïté, est un graphe (plus précisément un arbre) orienté qu'on définit par induction sur k_a :

- si $k_a = 1$, alors a est une arête conclusion d'un lien axiome ou affaiblissement (ce sont les seuls liens sans prémisses) de R , et G_a^R est le graphe sans liens constitué de la seule arête a . On écrira $G_a^R = \{a\}$.

Autrement, soit n le lien de R dont a est conclusion :

- si n est un lien avec deux prémisses et une conclusion (c.à.d. si $n = \otimes, \wp, \&, coad$), alors soient a_1 la prémisses gauche et a_2 la prémisses droite de n dans R . On a $k_{a_1} < k_a$ et $k_{a_2} < k_a$. G_a^R est le graphe obtenu en connectant les deux graphes $G_{a_1}^R$ et $G_{a_2}^R$ par le lien n de prémisses gauche a_1 , de prémisses droite a_2 , et de conclusion a . On écrira $G_a^R = G_{a_1}^R \cup G_{a_2}^R \cup \{n\} \cup \{a\}$.

- si n est un lien avec une prémisses et une conclusion (c.à.d. si $n = !, ?de, pax, \oplus$), alors soit a' la prémisses de n . On a $k_{a'} < k_a$. G_a^R est le graphe obtenu à partir de $G_{a'}^R$ en rajoutant le lien n de prémisses a' et de conclusion a . On écrira $G_a^R = G_{a'}^R \cup \{n\} \cup \{a\}$.

- si n est un lien $?co$ à $h \geq 2$ prémisses, alors soient (de gauche à droite) $a_1, \dots, a_h$ les prémisses de n dans R . On a $k_{a_i} < k_a$ (pour $i \in \{1, \dots, h\}$). G_a^R est le graphe obtenu en connectant les h graphes $G_{a_1}^R, \dots, G_{a_h}^R$ par le lien n de prémisses (de gauche à droite) $a_1, \dots, a_h$, et de conclusion a .

Le graphe partiel de a dans R (noté GP_a^R) est obtenu à partir de G_a^R en effaçant tous les liens pax ainsi que leurs conclusions.

7.2.4. REMARQUE. (et Définition). (i) Les graphes G_a^R et GP_a^R que nous venons de définir sont des arbres orientés. Mais de même qu'un réseau est défini à l'ordre des prémisses des liens $?co$ près, (c.à.d. que c'est en fait une classe d'équivalence de graphes) nous allons considérer comme étant équivalents deux arbres G_a^R ou GP_a^R qui ne diffèrent que par l'ordre des prémisses des contractions.

Plus précisément, si c est une arête d'un arbre G_a^R ou GP_a^R qui est conclusion d'un lien contraction ayant parmi ses prémisses a_1 et a_2 , alors nous dirons que l'arbre obtenu en permutant les deux sous-arbres de conclusions a_1 et a_2 est équivalent à G_a^R ou GP_a^R . On notera toujours G_a^R et GP_a^R la classe d'équivalence des arbres G_a^R et GP_a^R .

(ii) Si a est une arête du réseau R et si e est une expérience de R , on notera $e|_{G_a^R}$ la restriction de e aux arêtes de G_a^R . **Attention !** Si a' est une arête du

réseau R' et si e' est une expérience de R' , on écrira $e|_{G_a^R} = e'|_{G_{a'}^{R'}}$ lorsque *il existe* un arbre A (resp. A') de la classe d'équivalence G_a^R (resp. $G_{a'}^{R'}$) tels que $e|_A = e'|_{A'}$. Nous utiliserons la même convention pour les graphes partiels.

- 7.2.5. REMARQUE.** (i) Soit R un réseau, a_1 une arête de R conclusion du lien pax n et a_2 l'arête de R conclusion du lien déréluction au-dessus de n . On a alors $GP_{a_1}^R = GP_{a_2}^R$.
(ii) Soit R un réseau, a une arête de type A de R et c une arête de type C de G_a^R (resp. GP_a^R). Alors C est une occurrence de sous-formule de A .
(iii) Soit R un réseau, a et c deux arêtes de R . Si $a \in G_c^R$ alors G_a^R est un sous-graphe de G_c^R .
(iv) Soit R un réseau, $a_1, \dots, a_n$ les arêtes conclusions de R . $SPLO(R)$ (resp. $SPPO(R)$) n'est autre que la juxtaposition des graphes (resp. des graphes partiels) $G_{a_1}^R, \dots, G_{a_n}^R$ (resp. $GP_{a_1}^R, \dots, GP_{a_n}^R$).
(v) Soient a et a' deux arêtes du réseau R . $G_a \cap G_{a'} = \emptyset$ si et seulement si $a \notin G_{a'}$ et $a' \notin G_a$. Il en est de même pour les graphes partiels.

7.3 L'injectivité pour MLL

Remarquons que toute expérience d'un réseau R de MLL associe à toute arête une unique étiquette : dans le cas multiplicatif, une expérience de R est un étiquetage des arêtes de R .

Par ailleurs, si R est un réseau de MLL , alors $SPO(R) = SPLO(R) = SPPO(R)$.

Le lemme suivant est une simple remarque : si on connaît le type de toutes les conclusions du réseau R de MLL , alors on peut reconstruire R "aux liens axiome près".

7.3.1. LEMME. Si R et R' sont deux réseaux de même conclusion de MLL , alors $SPO(R) = SPO(R')$.

Démonstration : Evidente. □

Le lemme précédent met bien au clair que la seule information que nous devons extraire de l'interprétation d'un réseau R de MLL pour pouvoir affirmer que la sémantique est injective, c'est l'ensemble des couples d'arêtes de type atomique de R qui sont conclusion des liens axiome de R .

7.3.2. LEMME. Soit R (resp. R') un réseau de *MLL* de conclusion Γ , et soit e (resp. e') une expérience de R (resp. R') de conclusion γ (resp. γ'). Si $\gamma = \gamma'$, alors $(SPO(R) = SPO(R'))$ et $e|_{SPO(R)} = e'|_{SPO(R')}$.

Démonstration : Presque aussi évidente que celle du lemme précédent.

Si x est l'unique étiquette qu'une expérience ϵ d'un réseau T de *MLL* associe à une arête quelconque c , alors il existe une unique façon d'étiqueter les arêtes de G_c^T de telle sorte que l'étiquette de c soit x . En d'autres termes : $\epsilon|_{G_c^T}$ est univoquement déterminé par $\epsilon(c)$.

Soit alors a une conclusion de R et soit a' la conclusion de R' correspondante (voir remarque 6.2.2). De $e(a) = e'(a')$ on déduit alors $e|_{G_a^R} = e'|_{G_{a'}^{R'}}$, et donc $e|_{SPO(R)} = e'|_{SPO(R')}$. $\square$

Voici venu le seul endroit qui recquière une (toute petite) dose d'ingéniosité. Il faut trouver l'expérience appropriée qui permette de “refermer” les liens axiome. Nous introduisons pour cela la notion de “expérience injective”, que nous généraliserons à *MELL* tout entier dans la suite.

7.3.3. DÉFINITION. Soit R un réseau de *MLL* et e une expérience de R . Nous dirons que e est une **expérience injective** lorsque $\forall a, a'$ arêtes de même type atomique de R telles que $a \neq a'$, on a $e(a) \neq e(a')$.

Nous sommes maintenant en mesure de prouver le “théorème” d'injectivité pour *MLL*.

7.3.4. THÉORÈME. (Injectivité pour *MLL*). Soient R et R' deux réseaux de même conclusion. Si $[R] = [R']$, alors $R = R'$.

Démonstration : Soit e une expérience injective de R de conclusion γ .

Remarquons qu'une telle expérience existe toujours, parce qu'il n'y a aucune contrainte sur les étiquettes des arêtes d'un réseau standard de *MLL* : il n'y a aucun lien coupure (et ceci serait déjà suffisant pour la sémantique relationnelle), aucun lien ?co et aucun lien pax.

Puisque $[R] = [R']$, il existe une expérience e' de R' de conclusion γ . Par le lemme 7.3.2, on a alors $(SPO(R) = SPO(R'))$ et $e|_{SPO(R)} = e'|_{SPO(R')}$. Puisque e est injective, pour toute arête atomique a_1 de $SPO(R) = SPO(R')$ de type X , il existe une unique arête a_2 de $SPO(R) = SPO(R')$ de type $X^\perp$ telle que $e(a_1) = e'(a_1) = e(a_2) = e'(a_2)$. Donc $R = R'$. $\square$

7.3.5. REMARQUE. La démonstration ci-dessus met en évidence le phénomène suivant : toute expérience injective d'un réseau de *MLL* permet, à elle seule, de reconstituer entièrement le réseau. Ce qui peut conduire à se poser la question suivante : “à quoi servent les autres expériences ?”

Il ne faut pas oublier que la sémantique denotationnelle ne fournit pas une représentation *statique* d'une preuve π donnée : elle rend compte surtout de toutes les interactions possibles de π avec les autres démonstrations. Remarquons à ce propos que l'interprétation d'un réseau par l'ensemble des résultats de toutes ses expériences injectives n'est pas correcte : elle ne fournit pas un invariant de l'élimination des coupures (penser tout simplement à deux axiomes de même conclusion coupés entre eux).

Chapitre 8

Les expériences obsessionnelles

Pour prouver la conjecture 7.1.3, la première idée est celle d'appliquer exactement la même méthode qu'on a appliquée pour démontrer le théorème 7.3.4, c.à.d. (essentiellement) démontrer l'analogue du lemme 7.3.2 et généraliser aux réseaux de *MELL* la notion d'expérience injective.

Mais on se heurte immédiatement à une difficulté majeure : le type A d'une conclusion a d'un réseau R de *MELL* est loin d'être suffisant pour nous permettre de reconstituer le graphe G_a^R (à cause de la présence des liens $?co$ et pax , dont les prémisses et la conclusion ont toutes le même type). Bien pire, supposons qu'on connaisse les étiquettes que toutes les expériences de R associent à une arête a donnée qui est conclusion d'un lien n , supposons que cela nous suffise pour "deviner" quel est ce lien n (il faut bien l'espérer!), et supposons enfin qu'on découvre que le lien n en question est un lien $?co$ d'arité k : comment "deviner" pour une étiquette donnée de a la distribution des étiquettes sur les k prémisses de n ?

Notre idée fut celle de chercher, parmi toutes les expériences d'un réseau, des expériences suffisamment régulières, et de prouver pour ces expériences l'analogue du lemme 7.3.2. Les expériences obsessionnelles jouent précisément ce rôle : c'est ce que montre le corollaire 8.3.9 (et dans une moindre mesure le corollaire 8.2.3).

8.1 Définition et propriétés

8.1.1. DÉFINITION. (n -expérience obsessionnelle) Soit R un réseau, e une expérience de R , et $n \geq 1$ un entier. On dira que e est une n -expérience

obsessionnelle de R ssi

- (1) pour toute arête a de R de type atomique X , pour tout $x, y \in |\mathcal{X}|$, si $x \in e(a)$ et $y \in e(a)$, alors $x = y$
- (2) pour toute arête c de type $!C$ de R , $e(c) \neq \emptyset$ et pour tout $y \in e(c)$, $\text{card}(y) = n$.

L'entier n est aussi appelé le degré d'obsessionnalité de e .

Lorsque $n = 1$ nous dirons tout simplement que e est une 1-expérience.

8.1.2. REMARQUE. (i) Soit e une n -expérience obsessionnelle du réseau R et a une arête de R de type atomique et de profondeur p_a dans R . e satisfait les hypothèses du lemme 6.3.4 avec $\forall i \in \{1, \dots, p_a\} \ n_i = n$. On en déduit donc que $e(a)$ est un multienemble de cardinalité n^{p_a} , contenant n^{p_a} occurrences du même élément de $|\mathcal{X}|$.

(ii) Dans la cas des 1-expériences, les conditions (1) et (2) de la définition ci-dessus sont équivalentes à la seule condition (2).

(iii) Pour comprendre ce qu'est une n -expérience obsessionnelle, il convient peut-être de se représenter les choses de la façon suivante : prenons comme point de départ les sous-réseaux ne contenant aucune boîte, et considérons une expérience quelconque de chacun de ces sous-réseaux. Tant qu'on ne rencontre aucune porte de boîte, la propagation des étiquettes du haut vers le bas est totalement déterministe. Lorsqu'on rencontre une porte de boîte on s'arrête, et on attend les copains (c.à.d. qu'on attend que les prémisses de toutes les portes de la boîte aient une étiquette). Lorsque tout le monde est arrivé, on a une expérience e pour le contenu R_B de chaque boîte B ne contenant aucune autre boîte. On prend alors n copies de la même expérience e : on obtient une expérience de B (la condition sur les étiquettes des arêtes conclusions des liens pax et pal étant forcément satisfaite dans le cas cohérent). On recommence. Le lemme 8.1.4 montre bien que la définition qui précède est équivalente au processus informel que nous venons de décrire. Le lecteur aura certainement compris à présent la terminologie choisie pour ce type d'expérience.

8.1.3. LEMME. Soit e une n -expérience obsessionnelle du réseau R et R_1 un sous-réseau de R . Si $e|_{R_1} = \{e_1, \dots, e_l\}$, alors $\forall i \in \{1, \dots, l\} \ e_i$ est une n -expérience obsessionnelle de R_1 .

Démonstration : Conséquence directe du lemme 6.3.2 : pour toute arête a de R_1 $e(a) = e_1(a) \cup \dots \cup e_l(a)$, donc si e_i n'est pas une n -expérience obsessionnelle de R_1 alors e n'est pas une n -expérience obsessionnelle de R .

Plus formellement, supposons par absurde qu'il existe $j \in \{1, \dots, l\}$ t.q. e_j n'est pas une n -expérience obsessionnelle de R_1 .

Si e_j ne vérifie pas la condition (2) de la définition de n -expérience obsessionnelle, alors il existe une arête c de type $!C$ de R_1 t.q. une des deux conditions suivantes est satisfaite :

- (i) $e_j(c) = \emptyset$
- (ii) il existe $y \in e_j(c)$ t.q. $\text{card}(y) \neq n$

Dans le cas (i) il existe forcément une arête d de R_1 de type $!D$ t.q. $\emptyset \in e_j(d)$. Mais alors par le lemme 6.3.2 $\emptyset \in e(d)$, ce qui est contraire à la définition de n -expérience obsessionnelle.

Dans le cas (ii), par le lemme 6.3.2, $y \in e_j(c) \subseteq e(c)$, ce qui est contraire à la définition de n -expérience obsessionnelle.

e_j vérifie donc la condition (2) de la définition de n -expérience obsessionnelle. Soit a une arête atomique de type X de R_1 . Toujours par le lemme 6.3.2, si il existe $\alpha_1, \alpha_2 \in |\mathcal{X}|$ t.q. $\alpha_1 \neq \alpha_2$ et $\alpha_1, \alpha_2 \in e_j(a)$, alors $\alpha_1, \alpha_2 \in e(a)$ et e n'est pas une n -expérience obsessionnelle. $\square$

8.1.4. LEMME. Soient e et e' deux n -expériences obsessionnelles du réseau R . Si pour toute arête a de type atomique de R $e(a) = e'(a)$, alors $e = e'$.

Démonstration : On procède par induction sur R .

Si R est un axiome, alors c'est un axiome atomique et le résultat est trivial. Autrement, supposons tout d'abord que R soit une boîte B avec comme conclusion (resp. prémisses) de sa pal l'arête c (resp. l'arête d) et comme conclusions (resp. prémisses) de ses pax les arêtes $c_1, \dots, c_k$ (resp. $d_1, \dots, d_k$). Soit R_1 le plus grand sous-réseau de R contenu dans B , et soient $e_1, \dots, e_n$ (resp. $e'_1, \dots, e'_n$) les n expériences de R_1 telles que $e|_{R_1} = \{e_1, \dots, e_n\}$ (resp. $e'|_{R_1} = \{e'_1, \dots, e'_n\}$). $\forall i \in \{1, \dots, n\}$ e_i (resp. e'_i) est une n -expérience obsessionnelle de R_1 par le lemme 8.1.3. Pour toute arête atomique a de type X et de profondeur p dans R_1 (donc de profondeur $p+1$ dans R), on a $e(a) = e_1(a) \cup \dots \cup e_n(a)$ (resp. $e'(a) = e'_1(a) \cup \dots \cup e'_n(a)$) où $e(a)$ (resp. $e'(a)$) est un multiensemble de cardinalité n^{p+1} contenant n^{p+1} occurrences de $x \in |\mathcal{X}|$ (resp. $x' \in |\mathcal{X}|$) et $\forall i \in \{1, \dots, n\}$ $e_i(a)$ (resp. $e'_i(a)$) est un multiensemble de cardinalité n^p contenant n^p occurrences de $x_i \in |\mathcal{X}|$ (resp. de $x'_i \in |\mathcal{X}|$). Donc $x_1 = \dots = x_n = x$ (resp. $x'_1 = \dots = x'_n = x'$). Mais par hypothèse $e(a) = e'(a)$, donc $x = x'$ et $e_1(a) = \dots = e_n(a) = e'_1(a) = \dots = e'_n(a)$. Par hypothèse de récurrence $e_1 = \dots = e_n = e'_1 = \dots = e'_n$. Par définition d'expérience on a $e(c) = \{\{y_1, \dots, y_n\}\}$ (resp. $e'(c) = \{\{y'_1, \dots, y'_n\}\}$), avec $\{y_i\} = e_i(d) = e'_i(d) = \{y'_i\} \forall i \in \{1, \dots, n\}$, et

$\forall j \in \{1, \dots, k\} \ e(c_j) = \{z_1^j \cup \dots \cup z_n^j\}$ (resp. $e'(c_j) = \{z_1'^j \cup \dots \cup z_n'^j\}$), avec $\{z_i^j\} = e_i(d_j) = e'_i(d_j) = \{z_i'^j\} \ \forall i \in \{1, \dots, n\}$. Donc $e = e'$.

Si R n'est pas une boîte, soit n un dernier lien de R qui n'est pas porte d'une boîte. n est un lien affaiblissement, contraction, déréluction, ou bien un lien logique $\wp$ ou $\otimes$.

Si $n \neq \otimes$, alors soit R_1 le sous-réseau de R obtenu en éliminant le lien terminal n ainsi que sa conclusion. On aura $e|_{R_1} = \{e_1\}$ (resp. $e'|_{R_1} = \{e'_1\}$), et pour toute arête a de type atomique de R_1 (donc de R) $e_1(a) = e'_1(a)$. Par hypothèse de récurrence $e_1 = e'_1$, et donc $e = e'$.

Si tous les derniers liens de R qui ne sont pas des portes de boîtes sont des liens $\otimes$, alors il existe parmi eux un lien n qui est "scindant" (on utilise ici le théorème du $\otimes$ scindant de [Girard 87]) : en effaçant le lien n et sa conclusion, on obtient un graphe qui est l'union disjointe de deux sous-réseaux R_1 et R_2 de R . Pour le lien n en question, on aura $e|_{R_1} = \{e_1\}$ et $e|_{R_2} = \{e_2\}$ (resp. $e'|_{R_1} = \{e'_1\}$ et $e'|_{R_2} = \{e'_2\}$), et pour toute arête a de type atomique de R_1 (resp. de R_2), donc de R , $e_1(a) = e'_1(a)$ (resp. $e_2(a) = e'_2(a)$). Par hypothèse de récurrence on en déduira $e_1 = e'_1$ et $e_2 = e'_2$, et donc $e = e'$. $\square$

8.1.5. REMARQUE. Le lemme précédent est faux si on omet l'hypothèse d'obsessionnalité pour e et e' : considérons le réseau obtenu à partir d'un axiome de conclusion l'arête a_1 de type atomique X et l'arête a_2 de type $X^\perp$, en rajoutant ce qu'il faut (c.à.d. deux liens !, un lien $?de$ et deux liens pax , tous avec leurs conclusions) pour obtenir le réseau R de conclusion $!!X, ?X^\perp$. Les arêtes a_1 et a_2 sont contenues dans deux boîtes : appelons B_1 la plus petite et B_2 la plus grande. Soit e l'expérience du sous-réseau de R constitué par le seul lien axiome (et ses conclusions) telle que $e(a_1) = e(a_2) = x \in |\mathcal{X}|$. Soit e_1 (resp. e_2) l'expérience de R obtenue à partir de e en prenant 2 (resp. 4) copies de e pour sortir de B_1 et 8 (resp. 4) copies de l'expérience ainsi obtenue pour sortir de B_2 . Puisque $4 \times 4 = 2 \times 8$, on aura bien $e_1(a_1) = e_2(a_1)$ et $e_1(a_2) = e_2(a_2)$, et pourtant $e_1 \neq e_2$.

8.1.6. PROPOSITION. Soit e une n -expérience obsessionnelle du réseau R' et R un sous-réseau de R' . Si la profondeur des conclusions de R dans R' est p , alors $e|_R = \{e_1, \dots, e_{n^p}\}$, où $e_1 = \dots = e_{n^p}$ est une n -expérience obsessionnelle de R .

En particulier, pour toute arête a de type A de R' de profondeur p , $e(a)$ est un multiensemble de cardinalité n^p contenant n^p occurrences du même élément de $|\mathcal{A}|$.

Démonstration : Remarquons tout d'abord que la deuxième partie de la proposition est conséquence immédiate de la première : on prendra R sous-réseau de R' ayant a parmi ses conclusions et on utilisera l'égalité $e(a) = e_1(a) \cup \dots \cup e_n(a)$ (que fournit le lemme 6.3.2).

Pour prouver la première partie de la proposition, procédons par induction sur p .

Si $p = 0$, alors $e|_R = \{e_1\}$, et le résultat est évident.

Si les conclusions de R sont à profondeur $p + 1$ dans R' , alors soit B la boîte de profondeur 0 de R' contenant R . Soient $e_1, \dots, e_n$ les n expériences de R_B telles que $e|_{R_B} = \{e_1, \dots, e_n\}$. On sait que $\forall i \in \{1, \dots, n\}$ e_i est une n -expérience obsessionnelle de R_B (lemme 8.1.3). Par hypothèse de récurrence $\forall i \in \{1, \dots, n\}$, $e_i|_R = \{e_i^1, \dots, e_i^{n^p}\}$, avec $e_i^1 = \dots = e_i^{n^p}$ n -expérience obsessionnelle de R . Par ailleurs, $e|_R = e_1|_R \cup \dots \cup e_n|_R$ (définition 6.3.1). Pour conclure, il suffira donc de prouver que $e_1 = \dots = e_n$.

Pour toute arête a de type atomique de R (donc de R'), $e(a) = e_1(a) \cup \dots \cup e_n(a)$ et (par la remarque suivant la définition de n -expérience obsessionnelle) $e(a)$ est un multiensemble de cardinalité n^{p+1} , contenant n^{p+1} occurrences de $x \in |\mathcal{X}|$. Mais $\forall i \in \{1, \dots, n\}$ $e_i(a)$ est un multiensemble de cardinalité n^p , contenant n^p occurrences de $x_i \in |\mathcal{X}|$. Par conséquent $x_1 = \dots = x_n = x$.

$\forall i, j \in \{1, \dots, n\}$, e_i et e_j sont donc deux n -expériences obsessionnelles de R telles que pour toute arête a de type atomique de R , $e_i(a) = e_j(a)$. Par le lemme 8.1.4, on en déduit bien que $e_1 = \dots = e_n$. $\square$

8.1.7. REMARQUE. (i) Soit e une n -expérience obsessionnelle du réseau R , c une arête de type $!C$ de R et c' la prémisses du lien pal dont c est conclusion. Soit y l'unique élément de $e(c)$. Par définition d'expérience il existe $z_1, \dots, z_n \in e(c')$ t.q. $y = \{z_1, \dots, z_n\}$. Mais puisque (toujours par la proposition précédente) $e(c')$ contient un unique élément $z \in |\mathcal{C}|$, on aura $z_1 = \dots = z_n = z$.

(ii) Soit e (resp. e') une n -expérience obsessionnelle du réseau R (resp. du réseau R') et soit a (resp. a') une arête de profondeur p dans R (resp. de profondeur p' dans R'). La proposition précédente permet d'affirmer que si l'unique élément de $e(a)$ est égal à l'unique élément de $e'(a')$ et si $p = p'$, alors $e(a) = e'(a')$.

8.1.8. REMARQUE. Cette remarque ne concerne que la sémantique cohérente. Soit R un réseau sans liens ?co. Si on associe à toute arête de type atomique α un élément x_α de la trame de l'espace cohérent $\mathcal{A}$, alors il existe une (unique)

n -expérience obsessionnelle e_n de R t.q. l'unique élément de $e_n(\alpha)$ est x_α (pour toute arête α de type atomique). L'obsessionnalité de e_n nous garantit que les conditions de compatibilité requises par la définition 6.2.1 lors du passage dans les liens pax et pal sont satisfaites d'office. Remarquons que ce n'est pas le cas pour le passage dans les liens $?co$, et nous en discuterons de façon détaillée dans le chapitre 13.

8.2 Reconstruction des graphes partiels

Convention : A toute arête a de R de type A , une expérience e associe le multiensemble $e(a)$ des éléments de $|\mathcal{A}|$ qui sont étiquettes de a . Dans le cas particulier où e est une 1-expérience, le lecteur n'aura pas de mal à se convaincre que pour toute arête a de R $e(a)$ est un singleton. On utilisera la notation $e(a)$ tant pour indiquer le multiensemble $e(a)$ que pour indiquer l'unique élément qu'il contient (c.à.d. l'unique étiquette associée par l'expérience e à l'arête a) : c'est un peu abusif, mais beaucoup moins lourd. On dira aussi qu'une 1-expérience e d'un réseau R est une expérience pax-invisible, puisque si a (resp. a') est la prémisse (resp. la conclusion) d'un lien porte auxiliaire de R , alors $e(a) = e(a')$.

8.2.1. LEMME. Soit R un réseau, e une expérience pax-invisible de R et a une arête de R de type $?A$.

- (i) a est conclusion d'un lien affaiblissement ssi $e(a) = \emptyset$
- (ii) a est conclusion d'un lien dérélction ou porte auxiliaire ssi $e(a)$ est un singleton
- (iii) a est conclusion d'un lien contraction d'arité k ($k \geq 2$) ssi $e(a)$ est un élément de $|\mathcal{A}|$ de cardinalité k .

Démonstration :

- (i.1) Si a est conclusion d'un lien affaiblissement, alors par définition d'expérience $e(a) = \emptyset$.
- (ii.1) Soit a conclusion d'un lien dérélction ou porte auxiliaire n de R , et soit m le lien dérélction de R au-dessus de n (si n est un lien dérélction, on prendra $m = n$). Soit a_1 (resp. a_2) la prémisse (resp. la conclusion) de m . Puisque e est une expérience pax-invisible, $e(a) = e(a_2)$. Or, $e(a_2) = \{e(a_1)\}$ est un singleton par définition d'expérience.
- (i.2) Supposons que $e(a) = \emptyset$. Si a était conclusion d'un lien porte auxiliaire ou dérélction, alors par (ii.1) $e(a)$ serait un singleton. Si a était conclusion d'un lien contraction n d'arité $k \geq 2$, alors par le choix canonique de

l'emplacement des règles structurelles les prémisses de n seraient toutes des conclusions de liens dérélction ou porte auxiliaire, et donc toujours par (ii.1) l'étiquette associée par e à chacune de ces prémisses serait un singleton. Mais alors $e(a)$ serait un élément de $|\mathcal{A}|$ de cardinalité $k \geq 2$.

Donc a est forcément conclusion d'un lien affaiblissement.

- (iii.1) Si a est conclusion d'un lien contraction n d'arité k , alors on sait par le choix canonique de l'emplacement des règles structurelles que toutes les prémisses de n sont conclusions d'un lien dérélction ou bien d'un lien porte auxiliaire. Par (ii.1), on sait que l'étiquette de ces arêtes sera forcément un singleton, et puisque la sémantique est multiensembliste on peut en déduire que $e(a)$ est une clique de cardinalité k de $\mathcal{A}^\perp$.
- (ii.2) Si $e(a)$ est un singleton, alors a n'est pas conclusion d'un lien affaiblissement (par (i)) ni d'un lien contraction (par (iii.1)).
- (iii.2) Si $e(a)$ est un élément de cardinalité $k \geq 2$ de $|\mathcal{A}|$, alors par (i) et (ii) a ne peut être conclusion ni d'un lien affaiblissement, ni d'un lien dérélction, ni d'un lien porte auxiliaire, c.à.d. que a est conclusion d'un lien contraction n d'arité h . L'emplacement des règles structurelles dans un réseau standard, implique que les prémisses de n sont toutes conclusions d'un lien dérélction ou bien d'un lien porte auxiliaire. Mais alors par (ii) toutes ces arêtes sont étiquetées par des singletons, et $e(a)$ est un élément de $|\mathcal{A}|$ de cardinalité h . D'où $h = k$.

□

8.2.2. PROPOSITION. Soit R (resp. R') un réseau, et e (resp. e') une 1-expérience de R (resp. de R'). Soit a (resp. a') une arête de R (resp. de R') de type A . Si $e(a) = e'(a')$, alors :

- (i) $GP_a^R = GP_{a'}^{R'}$
- (ii) $e|_{GP_a^R} = e'|_{GP_{a'}^{R'}}$.

Démonstration : On procède par induction sur $s(GP_a^R)$, le nombre de liens de GP_a^R .

Si $s(GP_a^R) = 0$, alors a est une arête de R conclusion d'un lien axiome, et donc A est une formule atomique. Par conséquent a' est aussi conclusion d'un axiome et $GP_a^R = GP_{a'}^{R'}$. (ii) suit trivialement de l'hypothèse $e(a) = e'(a')$.

Autrement, soit n le lien de R dont a est conclusion.

- Si $n = \otimes$ ou bien $n = \emptyset$, alors l'arête a' de R' est aussi conclusion d'un lien $n' = \otimes$ ou bien $n' = \emptyset$. Soient a_1 et a_2 (resp. a'_1 et a'_2) les prémisses de n dans R (resp. de n' dans R'). a_i et a'_i ($i = 1, 2$) sont de même type, et par définition $GP_a^R = GP_{a_1}^R \cup GP_{a_2}^R \cup \{n\} \cup \{a\}$ et $GP_{a'}^{R'} = GP_{a'_1}^{R'} \cup GP_{a'_2}^{R'} \cup \{n'\} \cup \{a'\}$. On

a $s(GP_{a_1}^R) < s(GP_a^R)$, $s(GP_{a_2}^R) < s(GP_a^R)$ et $e(a_1) = e'(a'_1)$, $e(a_2) = e'(a'_2)$. L'hypothèse d'induction donne $GP_{a_1}^R = GP_{a'_1}^{R'}$, $GP_{a_2}^R = GP_{a'_2}^{R'}$ et $e|_{GP_{a_1}^R} = e'|_{GP_{a'_1}^{R'}}$, $e|_{GP_{a_2}^R} = e'|_{GP_{a'_2}^{R'}}$. La conclusion est alors immédiate.

- Si n est un lien pal, alors l'arête a' de R' est aussi conclusion d'un lien pal n' . Soit a_1 (resp. a'_1) la prémisses de n dans R (resp. de n' dans R'). a_1 et a'_1 sont de même type, et par définition $GP_a^R = GP_{a_1}^R \cup \{n\} \cup \{a\}$ et $GP_{a'}^{R'} = GP_{a'_1}^{R'} \cup \{n'\} \cup \{a'\}$. On a $s(GP_{a_1}^R) < s(GP_a^R)$ et $e(a_1) = e'(a'_1)$ (puisque e et e' sont des 1-expériences). L'hypothèse d'induction donne alors $GP_{a_1}^R = GP_{a'_1}^{R'}$ et $e|_{GP_{a_1}^R} = e'|_{GP_{a'_1}^{R'}}$, ce qui permet de conclure.

- Si n est un lien déréluction de R , alors $e(a)$ est un singleton, donc par hypothèse $e'(a')$ l'est aussi, et alors par le lemme 8.2.1 l'arête a' de R' est conclusion d'un lien déréluction ou porte auxiliaire m' de R' . Si m' est un lien déréluction on pose $n' = m'$; autrement soit n' le lien déréluction de R' au-dessus de m' . Si on appelle a''_1 l'arête de R' conclusion de n' , on a $GP_{a''_1}^{R'} = GP_{a'_1}^{R'}$, et $e(a) = e'(a') = e'(a''_1)$. Soit a_1 (resp. a'_1) l'arête de R (resp. de R') prémisses de n (resp. de n'). a_1 et a'_1 sont de même type, et par définition $GP_a^R = GP_{a_1}^R \cup \{n\} \cup \{a\}$ et $GP_{a'}^{R'} = GP_{a'_1}^{R'} \cup \{n'\} \cup \{a''_1\}$. On a $s(GP_{a_1}^R) < s(GP_a^R)$ et $e(a_1) = e'(a'_1)$. L'hypothèse d'induction donne alors $GP_{a_1}^R = GP_{a'_1}^{R'}$ et $e|_{GP_{a_1}^R} = e'|_{GP_{a'_1}^{R'}}$, ce qui permet de conclure.

- Si n est un lien pax de R , alors soit m la déréluction de R au-dessus de n , a_1 l'arête de R prémisses de m et a_2 l'arête de R conclusion de m . On a $GP_a^R = GP_{a_2}^R$ et $e(a) = e(a_2)$. Puisque e est une 1-expérience de R , le lemme 8.2.1 permet d'affirmer que $e(a) = e'(a')$ est un singleton. Ce même lemme permet d'en déduire que a' est conclusion d'un lien déréluction ou porte auxiliaire n' de R' . Si n' est un lien déréluction on pose $m' = n'$; autrement soit m' le lien déréluction de R' au-dessus de n' . Si on appelle a'_2 (resp. a'_1) l'arête de R' conclusion (resp. prémisses) de m' , on a $GP_{a'_2}^{R'} = GP_{a'_1}^{R'}$, et $e(a_2) = e(a) = e'(a') = e'(a'_2)$. a , a_1 , a_2 , a' , a'_1 , et a'_2 sont toutes des arêtes du même type, et par définition on aura alors $GP_{a'}^{R'} = GP_{a'_2}^{R'} = GP_{a'_1}^{R'} \cup \{m'\} \cup \{a'_2\}$ et $GP_a^R = GP_{a_2}^R = GP_{a_1}^R \cup \{m\} \cup \{a_2\}$. Par ailleurs $e(a_1) = e'(a'_1)$ et $s(GP_{a_1}^R) < s(GP_a^R)$, ce qui permet d'appliquer l'hypothèse de récurrence. On obtiendra donc $GP_{a_1}^R = GP_{a'_1}^{R'}$ et $e|_{GP_{a_1}^R} = e'|_{GP_{a'_1}^{R'}}$, ce qui permet de conclure.

- Si n est un lien affaiblissement, alors $e(a) = e'(a') = \emptyset$ et par le lemme 8.2.1 l'arête a' de R' est aussi conclusion d'un lien affaiblissement de R' . La conclusion est alors immédiate.

- Si n est un lien contraction d'arité k , alors par le lemme 8.2.1 l'arête a' de R' est aussi conclusion d'un lien contraction n' de R' d'arité k . Toujours par ce même lemme, on sait que $e(a)$ et $e(a')$ sont des ensemble de cardinal k . La position des règles structurelles dans les réseaux standard permet d'affirmer que toutes les prémisses de R (resp. de R') sont conclusions de liens dérélction ou porte auxiliaire. Chacune de ces arêtes sera donc étiquetée par un singleton. On peut alors choisir d'énumérer les prémisses $a_1, \dots, a_k$ de n et les prémisses $a'_1, \dots, a'_k$ de n' de telle sorte que $\forall i \in \{1, \dots, k\} e(a_i) = e'(a'_i)$. a_i et a'_i ($i = 1, \dots, k$) sont de même type, et par définition $GP_a^R = GP_{a_1}^R \cup \dots \cup GP_{a_k}^R \cup \{n\} \cup \{a\}$ et $GP_{a'}^{R'} = GP_{a'_1}^R \cup \dots \cup GP_{a'_k}^R \cup \{n'\} \cup \{a'\}$. Par ailleurs, $\forall i \in \{1, \dots, k\} s(GP_{a_i}^R) < s(GP_a^R)$ et $e(a_i) = e'(a'_i)$, ce qui permet d'appliquer l'hypothèse de récurrence. On obtiendra donc $\forall i \in \{1, \dots, k\} GP_{a_i}^R = GP_{a'_i}^{R'}$ et $e|_{GP_{a_i}^R} = e'|_{GP_{a'_i}^{R'}}$, ce qui permet de conclure. $\square$

8.2.3. COROLLAIRE. Soient R et R' deux réseaux de conclusion Γ , soit e (resp. e') une expérience pax-invisible de R (resp. de R') de conclusion γ (resp. γ'). Si $\gamma = \gamma'$, alors $SPPO(R) = SPPO(R')$ et $e|_{SPPO(R)} = e'|_{SPPO(R')}$.

Démonstration : Pour toute (occurrence de) formule A de Γ , appelons a (resp. a') l'arête conclusion de R (resp. de R') de type A . $e(a) = e'(a')$, donc par la proposition précédente $GP_a^R = GP_{a'}^{R'}$. Il suffit alors de rappeler que pour tout réseau T , $SPPO(T)$ n'est autre que l'union disjointe des graphes partiels des arêtes conclusions de T . $\square$

8.3 Reconstruction des graphes

8.3.1. DÉFINITION. Soient R un réseau, a une arête de R conclusion du lien contraction n d'arité k , pax ou dérélction. Soient $a_1, \dots, a_k$ (resp. $a'_1, \dots, a'_k$) les conclusions (resp. les prémisses) des k liens dérélction au-dessus de a ($k \geq 1$).

Le **graphe complet** du lien contraction n dans R , noté GC_a^R est le sous-graphe de G_a^R obtenu en effaçant $G_{a'_i}^R$ pour tout $i \in \{1, \dots, k\}$.

On dira que l'arête a (resp. les arêtes $a_1, \dots, a_k$) est (resp. sont) la conclusion (resp. les hypothèses) de GC_a^R . Pour toute arête c de GC_a^R on dira que le nombre p des liens pax traversés par le chemin de GC_a^R issu de c et ayant a comme arête terminale est la profondeur de c dans GC_a^R .

Lorsque nous utiliserons la définition précédente, l'arête a sera presque toujours la conclusion d'un lien contraction, mais avoir défini GC_a même lorsque a est conclusion d'un lien pax ou ?de se révèlera utile dans la suite.

8.3.2. REMARQUE. Soit a une arête du réseau R conclusion du lien contraction n d'arité k et soient $a_1, \dots, a_k$ les conclusions des k liens déréluction $n_1, \dots, n_k$ au-dessus de a . Pour toute arête c (resp. pour tout lien m) de GC_a^R t.q. $c \notin \{a, a_1, \dots, a_k\}$ (resp. $m \notin \{n, n_1, \dots, n_k\}$), c est la conclusion d'un lien pax (resp. m est un lien pax).

Dans le lemme suivant, toutes les lettres latines symbolisent des entiers positifs ou nuls, et on travaille (comme toujours) avec des *multiensembles*.

8.3.3. LEMME. Soient $1 \leq l, m < n$. On a $\forall p_1, \dots, p_l, p'_1, \dots, p'_m$: l'égalité $n^{p_1} + \dots + n^{p_l} = n^{p'_1} + \dots + n^{p'_m}$ implique l'égalité $\{p_1, \dots, p_l\} = \{p'_1, \dots, p'_m\}$ (en particulier $l = m$).

Démonstration : Ce résultat est une conséquence de l'unicité de la décomposition de tout entier en base n .

□

Soit b une arête de GC_a^R (a étant l'arête conclusion d'un lien ?co du réseau R), et soit c l'hypothèse de GC_a^R telle que $c \in G_b^R$. Dans les trois lemmes suivants, nous appellerons profondeur de b dans GC_a^R le nombre de liens pax traversés par l'unique chemin de R issu de c et ayant b comme arête terminale.

8.3.4. LEMME. Soit h l'arité maximale des liens contraction du réseau R , soit n un entier t.q. $n > \max(1, h)$. Soit e une n -expérience obsessionnelle de R . Soit a une arête de R de type ?A et soit x l'unique élément de $e(a)$ (suivant la proposition 8.1.6). Alors :

- (i) a est conclusion d'un lien affaiblissement ssi $x = \emptyset$
- (ii) a est conclusion d'un lien déréluction ssi x est un singleton
- (iii) a est conclusion d'un lien pax ssi il existe un entier $p \geq 1$ t.q. $\text{card}(x) = n^p$
- (iv) a est conclusion d'un lien contraction d'arité k ssi il existe $p_1, \dots, p_k$ entiers naturels t.q. $k \geq 2$ et $\text{card}(x) = n^{p_1} + \dots + n^{p_k}$.

Démonstration :

(i.1) et (ii.1) Remarquons tout d'abord que e est définie sur toutes les arêtes de R . Si a est conclusion d'un lien dérélction (resp. d'un lien affaiblissement), alors par définition d'expérience il existe $y \in |\mathcal{A}|$ t.q. $\{y\} \in e(a)$ (resp. $\emptyset \in e(a)$). Donc $x = \{y\}$ (resp. $x = \emptyset$).

(i.2) Si $\emptyset \in e(a)$, alors par ce qui précède a ne peut pas être conclusion d'un lien dérélction.

Si a est conclusion d'un lien pax, alors appelons a' la conclusion du lien dérélction au-dessus de a . Par ce qui précède, l'unique élément de $e(a')$ est un singleton, et alors puisque e est une n -expérience obsessionnelle on ne peut pas avoir $\emptyset \in e(a)$.

Si a est conclusion d'un lien contraction d'arité $k \geq 2$ ayant comme prémisses $a_1, \dots, a_k$ et si $x \in e(a)$, alors par définition d'expérience $\forall i \in \{1, \dots, k\}$ il existe $x_i \in e(a_i)$ t.q. $x = x_1 \cup \dots \cup x_k$. L'emplacement des règles structurelles dans un réseau standard implique que $\forall i \in \{1, \dots, k\}$ a_i est conclusion d'un lien pax ou dérélction. Mais dans les deux cas, par ce qui précède, $x_i \neq \emptyset \forall i \in \{1, \dots, k\}$, et donc $x \neq \emptyset$.

(ii.2) a ne peut pas être conclusion d'un lien affaiblissement par (i).

Si a est conclusion d'un lien pax dont la prémisse est a' et l'unique élément de $e(a)$ est un singleton $\{y\}$ de $\mathcal{A}$, alors par définition d'expérience il existe $x_1, \dots, x_n \in e(a')$ t.q. $\forall i \in \{1, \dots, n\}$, $x_i \in |\mathcal{A}|$ et $\{y\} = x_1 \cup \dots \cup x_n$. L'emplacement des règles structurelles dans un réseau standard implique que a' est conclusion d'un lien dérélction ou bien d'un lien pax, et par la proposition 8.1.6 $x_1 = \dots = x_n = z$. Or par (i) $z \neq \emptyset$, et la seule possibilité serait donc que $n = 1$ et $z = \{y\}$. Mais par hypothèse $n > 1$.

Si a est conclusion d'un lien contraction d'arité $k \geq 2$ dont les prémisses sont $a_1, \dots, a_k$, et si l'unique élément de $e(a)$ est un singleton $\{y\}$, alors $\forall i \in \{1, \dots, k\}$ il existe $x_i \in e(a_i)$ t.q. $\{y\} = x_1 \cup \dots \cup x_k$. L'emplacement des règles structurelles dans un réseau standard implique que $\forall i \in \{1, \dots, k\}$ a_i est conclusion d'un lien pax ou dérélction. Donc par (i) $\forall i \in \{1, \dots, k\}$ $x_i \neq \emptyset$. La seule possibilité serait alors que $k = 1$ et $x_1 = \{y\}$, mais $k \geq 2$.

(iii.1) Si a est conclusion d'un lien pax, alors soit a' la conclusion du lien dérélction au-dessus de a , p le nombre de liens pax traversés par le chemin issu de a' ayant a comme arête terminale, et soit $\{y\}$ l'unique élément de $e(a')$. Par définition de n -expérience obsessionnelle l'unique élément de $e(a)$ est le multiensemble de cardinalité n^p , contenant n^p occurrences de $y \in |\mathcal{A}|$.

(iv.1) Si a est conclusion d'un lien contraction d'arité k , alors soient $a_1, \dots, a_k$ les arêtes conclusion des k liens dérélction au-dessus de a et $p_1, \dots, p_k$ leurs profondeurs respectives dans GC_a^R . $\forall i \in \{1, \dots, k\}$, soit $\{y_i\}$ l'unique

élément de $e(a_i)$. Par définition de n -expérience obsessionnelle l'unique élément de $e(a)$ est le multiensemble de cardinalité $n^{p_1} + \dots + n^{p_k}$ contenant n^{p_i} occurrences de $y_i \forall i \in \{1, \dots, k\}$.

(iii.2) Si $\text{card}(x) > 1$, alors par (i) et (ii) a est conclusion d'un lien pax ou bien d'un lien contraction.

Si a est conclusion d'un lien contraction d'arité k avec $2 \leq k \leq h < n$, alors par ce qui précède il existe $p_1, \dots, p_k$ entiers naturels t.q. $\text{card}(x) = n^{p_1} + \dots + n^{p_k}$. Par le lemme 8.3.3 $\forall p \geq 1 \text{ card}(x) \neq n^p$. Donc si $\text{card}(x) = n^p$ avec $p \geq 1$, alors a est forcément conclusion d'un lien pax.

(iv.2) Puisque $k \geq 2$, $\text{card}(x) > 1$. Donc par (i) et (ii) a est conclusion d'un lien pax ou bien d'un lien contraction. Si a est conclusion d'un lien pax, alors par (iii) il existe $p \geq 1$ t.q. $\text{card}(x) = n^p$. Mais alors on aurait $n^p = n^{p_1} + \dots + n^{p_k}$, contredisant ainsi le lemme 8.3.3. a est donc bien conclusion d'un lien contraction. Soit l l'arité de ce lien. Par (iv.1), il existe $p'_1, \dots, p'_l$ entiers naturels t.q. $\text{card}(x) = n^{p'_1} + \dots + n^{p'_l}$. D'où $n^{p_1} + \dots + n^{p_k} = n^{p'_1} + \dots + n^{p'_l}$. Puisque $n > h \geq k, l \geq 1$ on peut appliquer le lemme 8.3.3 : on a donc $\{p'_1, \dots, p'_l\} = \{p_1, \dots, p_k\}$. En particulier $l = k$, c.à.d. que a est conclusion d'un lien contraction d'arité k . □

8.3.5. LEMME. Soit e (resp. e') une n -expérience obsessionnelle ($n > 1$) du réseau R (resp. R') et a (resp. a') la conclusion d'un lien contraction n de R (resp. n' de R') d'arité k (resp. k'). Soit c (resp. c') une des hypothèses de GC_a^R (resp. de $GC_{a'}^{R'}$) et soit $b \neq a$ (resp. $b' \neq a'$) une arête de GC_a^R (resp. de $GC_{a'}^{R'}$) telle que $c \in G_b^R$ (resp. $c' \in G_{b'}^{R'}$). Supposons que la profondeur de b dans GC_a^R soit égale à celle de b' dans $GC_{a'}^{R'}$.

Si $e(c) = e'(c')$, alors $e(b) = e'(b')$.

Démonstration : e et e' étant des n -expériences obsessionnelles avec $n > 1$, il existe un unique entier $h \geq 0$ tel que $\text{card}(e(c)) = \text{card}(e'(c')) = n^h$, donc l'arête c de R a la même profondeur que l'arête c' de R' .

Puisque l'arête b de GC_a^R a la même profondeur que l'arête b' de $GC_{a'}^{R'}$, le chemin issu de c ayant b comme arête terminale est égal au chemin issu de c' ayant b' comme arête terminale : il s'agit d'un certain nombre de liens pax avec leurs conclusions. e et e' étant des n -expériences obsessionnelles, de $e(c) = e'(c')$ on déduit $e(b) = e'(b')$. □

Soit e (resp. e') une n -expérience obsessionnelle du réseau R (resp. R'), avec $n > 1$. Soit a (resp. a') une arête de R (resp. de R') de type ?C conclusion d'un lien contraction d'arité $k \geq 2$. Soient $a_1, \dots, a_k$ (resp. $a'_1, \dots, a'_k$) les

hypothèses de GC_a^R (resp. de $GC_{a'}^{R'}$) et $p_1, \dots, p_k$ (resp. $p'_1, \dots, p'_k$) leurs profondeurs respectives dans GC_a^R (resp. dans $GC_{a'}^{R'}$).

Supposons que $e(a) = e'(a')$ et soit $t \in |\mathcal{C}|$ l'unique élément de $e(a) = e'(a')$ (suivant la proposition 8.1.6). $\forall x \in t$, on définit :

- $k_x := \text{card}(\{i \in \{1, \dots, k\} : \{x\} \text{ est l'unique élément de } e(a_i)\})$ (resp. $k'_x := \text{card}(\{i \in \{1, \dots, k\} : \{x\} \text{ est l'unique élément de } e'(a'_i)\})$)
- le multiensemble $\{p_1^x, \dots, p_{k_x}^x\} \subseteq \{p_1, \dots, p_k\}$ (resp. le multiensemble $\{p_1^{t_x}, \dots, p_{k'_x}^{t_x}\} \subseteq \{p'_1, \dots, p'_k\}$) des profondeurs dans GC_a^R (resp. dans $GC_{a'}^{R'}$) des arêtes $c \in \{a_1, \dots, a_k\}$ (resp. des arêtes $c' \in \{a'_1, \dots, a'_k\}$) telles que $\{x\}$ est l'unique élément de $e(c)$ (resp. de $e'(c')$).

8.3.6. LEMME. Si $\forall x \in t \ k_x = k'_x$ et $\{p_1^x, \dots, p_{k_x}^x\} = \{p_1^{t_x}, \dots, p_{k'_x}^{t_x}\}$, alors $GC_a^R = GC_{a'}^{R'}$ et $e|_{GC_a^R} = e'|_{GC_{a'}^{R'}}$.

Démonstration : On va prouver :

(*) $\forall c_1, c_2 \in \{a_1, \dots, a_k\}$, $c_1 \neq c_2$, c_1 (resp. c_2) de profondeur h_1 (resp. h_2) dans GC_a^R et t.q. $\{y_1\}$ (resp. $\{y_2\}$) est l'unique élément de $e(c_1)$ (resp. de $e(c_2)$), il existe une arête $c'_1 \in \{a'_1, \dots, a'_k\}$ (resp. une arête $c'_2 \in \{a'_1, \dots, a'_k\}$) de profondeur h_1 (resp. h_2) dans $GC_{a'}^{R'}$ t.q. $c'_1 \neq c'_2$ et $\{y_1\}$ (resp. $\{y_2\}$) est l'unique élément de $e'(c'_1)$ (resp. de $e'(c'_2)$).

Une fois (*) prouvée, on sait établir une injection (donc une bijection) f entre $\{a_1, \dots, a_k\}$ et $\{a'_1, \dots, a'_k\}$ telle que $\forall i \in \{1, \dots, k\}$ la profondeur de a_i dans GC_a^R est égale à celle de $f(a_i)$ dans $GC_{a'}^{R'}$, et l'unique élément de $e(a_i)$ est égal à l'unique élément de $e'(f(a_i))$. Par ailleurs, e et e' sont des n -expériences obsessionnelles avec $n > 1$, donc il existe un unique entier $h \geq 0$ tel que $\text{card}(e(a)) = \text{card}(e'(a')) = n^h$, c.à.d. que l'arête a de R a même profondeur que l'arête a' de R' . Donc $\forall i \in \{1, \dots, k\}$ l'arête a_i de R a même profondeur que l'arête $f(a_i)$ de R' . Mais alors, $\forall i \in \{1, \dots, k\}$ $e(a_i) = e'(f(a_i))$. Le lemme précédent (et l'hypothèse $e(a) = e'(a')$) permet(tent) alors de conclure.

Il reste donc à prouver (*) : soient $c_1, c_2 \in \{a_1, \dots, a_k\}$, $c_1 \neq c_2$, c_1 (resp. c_2) de profondeur h_1 (resp. h_2) dans GC_a^R et soit $\{y_1\}$ (resp. $\{y_2\}$) l'unique élément de $e(c_1)$ (resp. de $e(c_2)$).

Si $y_1 \neq y_2$, alors soient $\{p_1^{y_1}, \dots, p_{k_{y_1}}^{y_1}\}$ et $\{p_1^{y_2}, \dots, p_{k_{y_2}}^{y_2}\}$ (resp. $\{p_1^{y_1}, \dots, p_{k_{y_1}}^{y_1}\}$ et $\{p_1^{y_2}, \dots, p_{k_{y_2}}^{y_2}\}$) les multiensembles associés à y_1 (resp. y_2). Par hypothèse $h_1 \in \{p_1^{y_1}, \dots, p_{k_{y_1}}^{y_1}\} = \{p_1^{y_1}, \dots, p_{k_{y_1}}^{y_1}\}$ (resp. $h_2 \in \{p_1^{y_2}, \dots, p_{k_{y_2}}^{y_2}\} = \{p_1^{y_2}, \dots, p_{k_{y_2}}^{y_2}\}$), donc il existe $c'_1 \in \{a'_1, \dots, a'_k\}$ de profondeur h_1 (resp. $c'_2 \in \{a'_1, \dots, a'_k\}$ de

profondeur h_2) dans $GC_{a'}^{R'}$ telle que $\{y_1\}$ (resp. $\{y_2\}$) est l'unique élément de $e'(c'_1)$ (resp. de $e'(c'_2)$). Puisque $y_1 \neq y_2$, on a $c'_1 \neq c'_2$.

Si $y_1 = y_2 = y$, alors $\{p_1^{y_1}, \dots, p_{k_{y_1}}^{y_1}\} = \{p_1^{y_2}, \dots, p_{k_{y_2}}^{y_2}\} = \{p_1^y, \dots, p_{k_y}^y\}$ et il existe $1 \leq i \neq j \leq k_y$ t.q. $h_1 = p_i$ et $h_2 = p_j$. De $\{p_1^y, \dots, p_{k_y}^y\} = \{p_1^{ly}, \dots, p_{k_y}^{ly}\}$, on déduit alors l'existence de $c'_1, c'_2 \in \{a'_1, \dots, a'_k\}$, $c'_1 \neq c'_2$, c'_1 de profondeur p_i et c'_2 de profondeur p_j dans $GC_{a'}^{R'}$ t.q. $\{y\}$ est l'unique élément de $e'(c'_1)$ et de $e'(c'_2)$. $\square$

8.3.7. REMARQUE. Soient x et y des multiensembles. Si (pour $n \geq 1$) $x_1 \cup \dots \cup x_n = y_1 \cup \dots \cup y_n$, avec $x = x_1 = \dots = x_n$ et $y = y_1 = \dots = y_n$, alors $x = y$.

Démonstration : Il faut prouver que si $\alpha \in x$ et si $m_x(\alpha)$ (resp. $m_y(\alpha)$) est la multiplicité de α dans x (resp. dans y), alors $\alpha \in y$ et $m_x(\alpha) = m_y(\alpha)$. Si $\alpha \in x$, alors trivialement $\alpha \in y$. Par ailleurs l'hypothèse donne $n.m_x(\alpha) = n.m_y(\alpha)$, d'où $m_x(\alpha) = m_y(\alpha)$.

8.3.8. PROPOSITION. Soit a (resp. a') une arête de type A du réseau R (resp. du réseau R'). Soit h l'arité maximale des liens contraction de R et R' , et soit $n > \max(1, h)$.

Soit e (resp. e') une n -expérience obsessionnelle de R (resp. de R'). Si $e(a) = e'(a')$, alors :

- (i) $G_a^R = G_{a'}^{R'}$
- (ii) $e|_{G_a^R} = e'|_{G_{a'}^{R'}}$.

Démonstration : On procède par induction sur $s(G_a^R)$, le nombre de liens de G_a^R .

Si $s(G_a^R) = 0$, alors a est une arête de R conclusion d'un lien axiome, et donc A est une formule atomique. Par conséquent a' est aussi conclusion d'un axiome et $G_a^R = G_{a'}^{R'}$. (ii) suit trivialement de l'hypothèse $e(a) = e'(a')$. Autrement, soit m le lien de R dont a est conclusion.

- Si $m = \otimes$ ou bien $m = \oslash$, alors l'arête a' de R' est aussi conclusion d'un lien $m' = \otimes$ ou bien $m' = \oslash$. Soient a_1 et a_2 (resp. a'_1 et a'_2) les prémisses de m dans R (resp. de m' dans R'). a_i et a'_i ($i = 1, 2$) sont de même type, et par définition $G_a^R = G_{a_1}^R \cup G_{a_2}^R \cup \{m\} \cup \{a\}$ et $G_{a'}^{R'} = G_{a'_1}^{R'} \cup G_{a'_2}^{R'} \cup \{m'\} \cup \{a'\}$. On a $s(G_{a_1}^R) < s(G_a^R)$, $s(G_{a_2}^R) < s(G_a^R)$. Soit x l'unique élément de $e(a) = e'(a')$ (suivant la proposition 8.1.6). Par définition d'expérience, il existe $x_1 \in e(a_1)$ et $x_2 \in e(a_2)$ (resp. $x'_1 \in e'(a'_1)$ et $x'_2 \in e'(a'_2)$) t.q. $x = (x_1, x_2) = (x'_1, x'_2)$.

Donc $x_1 = x'_1$ est l'unique élément de $e(a_1)$ et de $e'(a'_1)$ et $x_2 = x'_2$ est l'unique élément de $e(a_2)$ et de $e'(a'_2)$. Puisque $e(a) = e'(a')$, les arêtes a et a' ont la même profondeur, donc a_1 , a_2 , a'_1 et a'_2 ont la même profondeur. Par conséquent $e(a_1) = e'(a'_1)$ et $e(a_2) = e'(a'_2)$. L'hypothèse d'induction donne alors $G_{a_1}^R = G_{a'_1}^{R'}$, $G_{a_2}^R = G_{a'_2}^{R'}$ et $e|_{G_{a_1}^R} = e'|_{G_{a'_1}^{R'}}$, $e|_{G_{a_2}^R} = e'|_{G_{a'_2}^{R'}}$. La conclusion est alors immédiate.

- Si m est un lien pal, alors l'arête a' de R' est aussi conclusion d'un lien pal m' . Soit a_1 (resp. a'_1) la prémisse de m dans R (resp. de m' dans R'). a_1 et a'_1 sont de même type, et par définition $G_a^R = G_{a_1}^R \cup \{m\} \cup \{a\}$ et $G_{a'}^R = G_{a'_1}^R \cup \{m'\} \cup \{a'\}$. On a $s(G_{a_1}^R) < s(G_a^R)$. Soit x l'unique élément de $e(a) = e'(a')$ (suivant la proposition 8.1.6). Par définition de n -expérience obsessionnelle, il existe $x_1, \dots, x_n \in e(a_1)$ (resp. $x'_1, \dots, x'_n \in e'(a'_1)$) t.q. $x = \{x_1, \dots, x_n\} = \{x'_1, \dots, x'_n\}$. La proposition 8.1.6 donne alors $x_1 = \dots = x_n = x'_1 = \dots = x'_n$ c.à.d. que l'unique élément de $e(a_1)$ est aussi l'unique élément de $e'(a'_1)$. Puisque $e(a) = e'(a')$, les arêtes a et a' ont la même profondeur, donc a_1 et a'_1 ont aussi la même profondeur. D'où $e(a_1) = e'(a'_1)$. L'hypothèse d'induction donne alors $G_{a_1}^R = G_{a'_1}^{R'}$ et $e|_{G_{a_1}^R} = e'|_{G_{a'_1}^{R'}}$, ce qui permet de conclure.

- Si m est un lien dérélction de R , alors $A = ?C$ et par le lemme 8.3.4, l'unique élément de $e(a) = e'(a')$ est un singleton. Par ce même lemme, a' est aussi conclusion d'un lien dérélction m' . Soit a_1 (resp. a'_1) la prémisse de m dans R (resp. de m' dans R'). a_1 et a'_1 sont de même type, et par définition $G_a^R = G_{a_1}^R \cup \{m\} \cup \{a\}$ et $G_{a'}^R = G_{a'_1}^R \cup \{m'\} \cup \{a'\}$. On a $s(G_{a_1}^R) < s(G_a^R)$. Soit x l'unique élément de $e(a) = e'(a')$ (suivant la proposition 8.1.6). Par définition d'expérience, il existe $y_1 \in e(a_1)$ (resp. $y'_1 \in e'(a'_1)$) t.q. $x = \{y_1\} = \{y'_1\}$. Donc $y_1 = y'_1$, c.à.d. que l'unique élément de $e(a_1)$ est aussi l'unique élément de $e'(a'_1)$. Puisque $e(a) = e'(a')$, les arêtes a et a' ont la même profondeur, donc a_1 et a'_1 ont aussi la même profondeur. D'où $e(a_1) = e'(a'_1)$. L'hypothèse d'induction donne alors $G_{a_1}^R = G_{a'_1}^{R'}$ et $e|_{G_{a_1}^R} = e'|_{G_{a'_1}^{R'}}$, ce qui permet de conclure.

- Si m est un lien pax de R , alors soit x l'unique élément de $e(a)$ (suivant la proposition 8.1.6). Par le lemme 8.3.4, il existe $p \geq 1$ t.q. $\text{card}(x) = n^p$. Par ce même lemme, puisque $e(a) = e'(a')$, a' est forcément conclusion d'un lien pax m' de R' . Soit a_1 (resp. a'_1) la prémisse de m dans R (resp. de m' dans R'). a_1 et a'_1 sont de même type, et par définition $G_a^R = G_{a_1}^R \cup \{m\} \cup \{a\}$ et $G_{a'}^R = G_{a'_1}^R \cup \{m'\} \cup \{a'\}$. On a $s(G_{a_1}^R) < s(G_a^R)$. Par définition de n -expérience obsessionnelle, il existe $x_1, \dots, x_n \in e(a_1)$ (resp. $x'_1, \dots, x'_n \in e'(a'_1)$) t.q. $x = x_1 \cup \dots \cup x_n = x'_1 \cup \dots \cup x'_n$. La proposition 8.1.6 et la remarque précédente

donnent alors $x_1 = \dots = x_n = x'_1 = \dots = x'_n$, c.à.d. que l'unique élément de $e(a_1)$ est aussi l'unique élément de $e'(a'_1)$. Puisque $e(a) = e'(a')$, les arêtes a et a' ont la même profondeur, donc a_1 et a'_1 ont aussi la même profondeur. D'où $e(a_1) = e'(a'_1)$. L'hypothèse d'induction donne alors $G_{a_1}^R = G_{a'_1}^{R'}$ et $e|_{G_{a_1}^R} = e'|_{G_{a'_1}^{R'}}$, ce qui permet de conclure.

- Si m est un lien affaiblissement, alors par le lemme 8.3.4, l'unique élément de $e(a) = e'(a')$ est $\emptyset$. Par ce même lemme, puisque $e(a) = e'(a')$, a' est forcément conclusion d'un lien affaiblissement et la conclusion est alors immédiate.

- Si m est un lien contraction d'arité $k \geq 2$, alors soit t l'unique élément de $e(a) = e'(a')$. Par le lemme 8.3.4, il existe $p_1, \dots, p_k$ entiers naturels t.q. $\text{card}(x) = n^{p_1} + \dots + n^{p_k}$. Par ce même lemme, puisque $e(a) = e'(a')$, a' est forcément conclusion d'un lien contraction m' d'arité k .

Soient $a_1, \dots, a_k$ (resp. $a'_1, \dots, a'_k$) les hypothèses de GC_a^R (resp. de $GC_{a'}^{R'}$) et $p_1, \dots, p_k$ (resp. $p'_1, \dots, p'_k$) leurs profondeurs respectives dans GC_a^R (resp. dans $GC_{a'}^{R'}$). Avec ces notations (qui sont celles du lemme 8.3.6), on va prouver que les hypothèses du lemme 8.3.6 sont satisfaites. Soit $x \in t$ et $\{p_1^x, \dots, p_{k_x}^x\} \subseteq \{p_1, \dots, p_k\}$ (resp. $\{p_1'^x, \dots, p_{k'_x}'^x\} \subseteq \{p'_1, \dots, p'_k\}$) le multi-ensemble des profondeurs dans GC_a^R (resp. dans $GC_{a'}^{R'}$) des arêtes $c \in \{a_1, \dots, a_k\}$ (resp. des arêtes $c' \in \{a'_1, \dots, a'_k\}$) telles que $\{x\}$ est l'unique élément de $e(c)$ (resp. de $e'(c')$). Soit α la cardinalité de t et soit α_x le nombre d'occurrences de x dans t . Puisque e et e' sont des n -expériences obsessionnelles, on a $\alpha_x = n^{p_1^x} + \dots + n^{p_{k_x}^x} = n^{p_1'^x} + \dots + n^{p_{k'_x}'^x}$. Puisque $n > h \geq k_x, k'_x \geq 1$, on peut appliquer le lemme 8.3.3, ce qui donne $\{p_1^x, \dots, p_{k_x}^x\} = \{p_1'^x, \dots, p_{k'_x}'^x\}$. Les hypothèses du lemme 8.3.6 sont donc satisfaites, et on a $GC_a^R = GC_{a'}^{R'}$ et $e|_{GC_a^R} = e'|_{GC_{a'}^{R'}}$. En particulier, en renommant éventuellement les hypothèses de GC_a^R et de $GC_{a'}^{R'}$ (rappelons au passage que G_a^R et $G_{a'}^{R'}$ sont des classes d'équivalences d'arbres), on a $\forall i \in \{1, \dots, k\} e(a_i) = e'(a'_i)$, et $s(G_{a_i}^R) < s(G_{a_i}^R)$. L'hypothèse d'induction donne alors $G_{a_i}^R = G_{a'_i}^{R'}$ et $e|_{G_{a_i}^R} = e'|_{G_{a'_i}^{R'}}$, $\forall i \in \{1, \dots, k\}$. D'où la conclusion.

□

8.3.9. COROLLAIRE. Soient R et R' deux réseaux de conclusion Γ , soit h l'arité maximale des liens contraction de R et R' , et soit $n > \max(1, h)$. Soit e (resp. e') une n -expérience obsessionnelle de R (resp. de R') de conclusion γ (resp. γ'). Si $\gamma = \gamma'$, alors $SPLO(R) = SPLO(R')$ et $e|_{SPLO(R)} = e'|_{SPLO(R')}$.

Démonstration : Pour toute (occurrence de) formule A de Γ , appelons a (resp. a') l'arête conclusion de R (resp. de R') de type A . $e(a) = e'(a')$, donc par la proposition précédente $G_a^R = G_{a'}^{R'}$. Il suffit alors de rappeler que pour tout réseau T , $SPLO(T)$ n'est autre que l'union disjointe des graphes des arêtes conclusions de T . $\square$

Le corollaire précédent est l'analogue (dans $MELL$) du lemme 7.3.2. Nous allons donc essayer de reproduire l'argument de la preuve du théorème 7.3.4 pour prouver la conjecture 7.1.3. Il faut alors généraliser la notion d'expérience injective à $MELL$. Nous introduisons par la même occasion la notion d'expérience simple dont nous ferons aussi usage dans la suite.

8.3.10. DÉFINITION. Soit R un réseau de $MELL$ et e une n -expérience obsessionnelle de R . Nous dirons que e est **injective** (resp. **simple**) lorsque $\forall a, a'$ arêtes de même type atomique de R telles que $a \neq a'$, l'unique élément de $e(a)$ est différent de (resp. égal à) l'unique élément de $e(a')$.

Tentons à présent de prouver la conjecture 7.1.3 : soit e une n -expérience obsessionnelle injective de R (**Question 1 : existe-t-elle ?**) de conclusion γ . Puisque $[R] = [R']$, il existe une expérience e' de R' de conclusion γ (**Question 2 : e' est-elle n -obsessionnelle ?**).

Si la réponse aux deux questions précédentes est positive, alors on peut appliquer le corollaire 8.3.9 : $SPLO(R) = SPLO(R')$ et $e|_{SPLO(R)} = e'|_{SPLO(R')}$. L'injectivité de e (donc de e') nous permet de conclure que dans ce cas $SPL(R) = SPL(R')$, et nous verrons que *en l'absence d'affaiblissements* ceci implique que $R = R'$.

Les chapitres suivants (qui réservent quelques surprises) fournissent des éléments de réponses aux deux questions ci-dessus.

Remarquons que si pour un fragment quelconque de LL (sans affaiblissements) la réponse aux deux questions ci-dessus était positive, une remarque analogue à 7.3.5 s'appliquerait (ce qui, en présence d'exponentielles, est plus frappant) : toute expérience n -obsessionnelle et injective d'un réseau de ce fragment permet à *elle toute seule* de reconstituer entièrement le réseau.

Chapitre 9

Résultats obsessionnels

Le but du présent chapitre est celui de donner des éléments de réponse à la question 2 de la fin du chapitre précédent. La question qui se pose est celle de savoir si, à partir du résultat d'une expérience donnée, on peut oui ou non inférer que celle-ci est obsessionnelle.

Nous montrons que la réponse est (essentiellement) positive dans le cas de la sémantique cohérente (proposition 9.3.5), et ce grâce à son “uniformité” (le lemme 9.3.1). Pour la sémantique relationnelle, nous prouvons que la réponse est positive pour les expériences obsessionnelles simples (proposition 9.2.2), et nous conjecturons qu'elle demeure (essentiellement) positive dans le cas général (conjecture 9.2.3).

9.1 Le cas général

9.1.1. LEMME. Soit R un réseau, c une arête de R de type $?C$ (resp. de type $!C$) conclusion du lien porte auxiliaire (resp. porte principale) n et soit c' la prémisses de n . Soit A une occurrence de sous-formule de C et e une expérience de R .

Soit $y \in e(c)$ t.q. $y = y_1 \cup \dots \cup y_k$ (resp. $y = \{y_1, \dots, y_k\}$) avec $y_i \in e(c')$ $\forall i \in \{1, \dots, k\}$. Alors $|y|_A = |y_1|_A \cup \dots \cup |y_k|_A$.

Démonstration : C'est une application des définitions. On procède par induction sur la complexité de $C \setminus A$.

Si $C = A$, alors dans le cas où c est de type $!C$ $|y|_A = \{x \in |A| : x \in y = \{y_1, \dots, y_k\}\} = \{y_1, \dots, y_k\}$. Par ailleurs, puisque c' est de type A et $y_i \in e(c')$, $|y_i|_A = \{y_i\} \forall i \in \{1, \dots, k\}$. Donc $|y|_A = |y_1|_A \cup \dots \cup |y_k|_A$.

Dans le cas où c est de type $?C$, $|y|_A = \{x \in |A| : x \in y = y_1 \cup \dots \cup y_k\} =$

$y_1 \cup \dots \cup y_k$. Par ailleurs, puisque c' est de type $?A$ et $y_i \in e(c')$, $|y_i|_A = \{x \in |\mathcal{A}| : x \in y_i\} = y_i \ \forall i \in \{1, \dots, k\}$. Donc $|y|_A = |y_1|_A \cup \dots \cup |y_k|_A$.

Autrement, on peut traiter les deux cas (c de type $!C$ et c de type $?C$) ensemble. Soit D la plus petite occurrence de sous-formule de C telle que A soit sous-formule stricte de D . On aura $D = A \otimes E, E \otimes A, A \wp E, E \wp A, ?A, !A$.

Si $D = A \otimes E$ ou bien $D = A \wp E$, alors par hypothèse de récurrence $|y|_D = |y_1|_D \cup \dots \cup |y_k|_D$. Par définition, on aura $|y|_A = \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y|_D\} = \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y_1|_D \cup \dots \cup |y_k|_D\} = \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y_1|_D\} \cup \dots \cup \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y_k|_D\} = |y_1|_A \cup \dots \cup |y_k|_A$.

Si $D = !A$ ou bien $D = ?A$, alors par hypothèse de récurrence $|y|_D = |y_1|_D \cup \dots \cup |y_k|_D$. Par définition, on aura $|y|_A = \{x \in |\mathcal{A}| : x \in z \in |y|_D\} = \{x \in |\mathcal{A}| : x \in z \in |y_1|_D \cup \dots \cup |y_k|_D\} = \{x \in |\mathcal{A}| : x \in z \in |y_1|_D\} \cup \dots \cup \{x \in |\mathcal{A}| : x \in z \in |y_k|_D\} = |y_1|_A \cup \dots \cup |y_k|_A$. $\square$

9.1.2. LEMME. Soit R un réseau, $c_1, \dots, c_k$ les k arêtes de R de type $?C$ prémisses du lien contraction n d'arité k et de conclusion c . Soit A une occurrence de sous-formule de C et e une expérience de R . Soit $y \in e(c)$ et soient $y_i \in e(c_i)$ ($i \in \{1, \dots, k\}$) tels que $y = y_1 \cup \dots \cup y_k$.

On a $|y|_A = |y_1|_A \cup \dots \cup |y_k|_A$.

Démonstration : Comme pour le lemme précédent, c'est une application des définitions. On procède par induction sur la complexité de $C \setminus A$.

Si $C = A$, alors $|y|_A = \{x \in |\mathcal{A}| : x \in y = y_1 \cup \dots \cup y_k\} = y_1 \cup \dots \cup y_k$. Par ailleurs, puisque $\forall i \in \{1, \dots, k\}$ c_i est de type $?A$ et $y_i \in e(c_i)$, on a $|y_i|_A = \{x \in |\mathcal{A}| : x \in y_i\} = y_i$. Donc $|y|_A = |y_1|_A \cup \dots \cup |y_k|_A$.

Autrement, soit D la plus petite occurrence de sous-formule de C telle que A soit sous-formule stricte de D . On aura $D = A \otimes E, E \otimes A, A \wp E, E \wp A, ?A, !A$.

Si $D = A \otimes E$ ou bien $D = A \wp E$, alors par hypothèse de récurrence $|y|_D = |y_1|_D \cup \dots \cup |y_k|_D$. Par définition, on aura $|y|_A = \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y|_D\} = \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y_1|_D \cup \dots \cup |y_k|_D\} = \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y_1|_D\} \cup \dots \cup \{x_1 \in |\mathcal{A}| : \exists x_2 \in |\mathcal{E}| t.q. (x_1, x_2) \in |y_k|_D\} = |y_1|_A \cup \dots \cup |y_k|_A$.

Si $D = !A$ ou bien $D = ?A$, alors par hypothèse de récurrence $|y|_D = |y_1|_D \cup \dots \cup |y_k|_D$. Par définition, on aura $|y|_A = \{x \in |\mathcal{A}| : x \in z \in |y|_D\} = \{x \in |\mathcal{A}| : x \in z \in |y_1|_D \cup \dots \cup |y_k|_D\} = \{x \in |\mathcal{A}| : x \in z \in |y_1|_D\} \cup \dots \cup \{x \in |\mathcal{A}| : x \in z \in |y_k|_D\} = |y_1|_A \cup \dots \cup |y_k|_A$. $\square$

9.1.3. LEMME. Soit R un réseau, c une arête de R de type $?C$ (resp. de type $C_1 \otimes C_2$ ou bien $C_1 \wp C_2$) conclusion du lien dérélction (resp. $\otimes$ ou bien $\wp$) n de prémisses c' (resp. de prémisses c_1 et c_2). Soit A une occurrence de sous-formule de C (resp. de C_i , pour un certain $i \in \{1, 2\}$) et e une expérience de R .

Soit $y \in e(c)$ et $y' \in e(c')$ (resp. $y_1 \in e(c_1)$ et $y_2 \in e(c_2)$) t.q. $y = \{y'\}$ (resp. $y = (y_1, y_2)$). Alors $|y|_A = |y'|_A$ (resp. $|y|_A = |y_i|_A$).

Démonstration : Comme pour les lemmes précédents, c'est une application des définitions. On procède par induction sur la complexité de $C \setminus A$ (resp. sur la complexité de $C_i \setminus A$).

Si $C = A$ (resp. si $C_i = A$), alors $|y|_A = \{x \in |\mathcal{A}| : x \in y\} = \{y'\} = |y'|_A$ (resp. $|y|_A = \{y_i\} = |y_i|_A$).

Autrement, soit D la plus petite occurrence de sous-formule de C telle que A soit sous-formule stricte de D . On aura $D = A \otimes E, E \otimes A, A \wp E, E \wp A, ?A, !A$. Le résultat découle alors d'une simple application de l'hypothèse de récurrence. $\square$

9.1.4. LEMME. Soit R un réseau, c une arête de R de type C et e une expérience de R . Pour tout x , les deux affirmations suivantes sont équivalentes :

- (1) Il existe une arête d de type A de G_c^R , avec $A \neq ?F$ pour toute formule F , telle que $x \in e(d)$
- (2) Il existe $y \in e(c)$ tel que $x \in |y|_A$, avec $A \neq ?F$ pour toute formule F .

Démonstration : On procède par induction sur $s(G_c)$, le nombre de liens du graphe G_c .

Si $s(G_c) = 0$, alors C est une formule atomique et le lemme est vrai avec $A = C$, $d = c$, $y = x$.

Autrement, soit n le lien de R de conclusion c . Si $C = A$, alors le lemme est vrai avec $c = d$ et $y = x$ (se souvenir que $A \neq ?F$ pour toute formule F). Supposons donc que A soit une sous-formule stricte de C . On aura alors forcément $c \neq d$.

- Si $n = \otimes$ ou bien $n = \wp$, alors $C = C_1 \otimes C_2$ ou bien $C = C_1 \wp C_2$. Soient c_1 et c_2 les prémisses de n , de type C_1 et C_2 respectivement. $s(G_{c_i}) < s(G_c)$ pour $i = 1, 2$.

Soit d une arête de type A de G_c^R telle que $x \in e(d)$. A est alors une occurrence de sous-formule de C , donc c'est une occurrence de sous-formule d'une et une seule des deux sous-formules C_1 et C_2 de C . Supposons par exemple que A soit sous-formule de C_1 , c.à.d. que $d \in G_{c_1}$. Par hypothèse de récurrence il existe $y_1 \in e(c_1)$ tel que $x \in |y_1|_A$. Soit alors $y_2 \in e(c_2)$ tel que

$y = (y_1, y_2) \in e(c)$. Puisque A est sous-formule de C_1 , par le lemme 9.1.3 $|y|_A = |y_1|_A \ni x$.

Réciproquement, soit $y \in e(c)$ tel que $x \in |y|_A$. Puisque $y \in e(c)$, on a forcément $y = (y_1, y_2)$ avec $y_i \in e(c_i)$ ($i = 1, 2$). Par ailleurs, de $|y|_A \neq \emptyset$ on déduit que A est une occurrence de sous-formule de C . Si on suppose cette fois que A est sous-formule de C_2 , on aura par le lemme 9.1.3 $|y|_A = |y_2|_A$, et donc $x \in |y_2|_A$. Par hypothèse de récurrence, on en déduit qu'il existe une arête d de type A de $G_{c_2}^R$ telle que $x \in e(d)$. Mais l'arête d de $G_{c_2}^R$ est aussi une arête de G_c^R .

- Si n est un lien contraction d'arité k , alors soient $c_1, \dots, c_k$ les k arêtes de R prémisses de n . $\forall i \in \{1, \dots, k\} s(G_{c_i}) < s(G_c)$.

Soit d une arête de type A de G_c telle que $x \in e(d)$. Il existe un unique $i \in \{1, \dots, k\}$ tel que d est arête de G_{c_i} . Par hypothèse de récurrence il existe $y_i \in e(c_i)$ tel que $x \in |y_i|_A$. Soit alors $\forall j \in \{1, \dots, k\}, j \neq i, y_j \in e(c_j)$ tel que $y = y_1 \cup \dots \cup y_i \cup \dots \cup y_k \in e(c)$. Par le lemme 9.1.2 $x \in |y|_A$.

Réciproquement, soit $y \in e(c)$ tel que $x \in |y|_A$. Puisque $y \in e(c)$, on a forcément $y = y_1 \cup \dots \cup y_k$, avec $y_i \in e(c_i) \forall i \in \{1, \dots, k\}$. Par le lemme 9.1.2, il existe $i \in \{1, \dots, k\}$ tel que $x \in |y_i|_A$. Par hypothèse de récurrence, on en déduit qu'il existe une arête d de type A de $G_{c_i}^R$ telle que $x \in e(d)$. Mais l'arête d de $G_{c_i}^R$ est aussi une arête de G_c^R .

- Si n est un lien affaiblissement, alors le résultat est trivial.

- Si n est un lien déréluction, alors soit c_1 la prémisse de n . $s(G_{c_1}) < s(G_c)$. Soit d une arête de type A de G_c telle que $x \in e(d)$. d est aussi une arête de G_{c_1} , et par hypothèse de récurrence, il existe $y_1 \in e(c_1)$ t.q. $x \in |y_1|_A$. Soit alors $y = \{y_1\} \in e(c)$. On a par le lemme 9.1.3 $|y|_A = |y_1|_A \ni x$.

Réciproquement, soit $y \in e(c)$ tel que $x \in |y|_A$. Il existe $y_1 \in e(c_1)$ t.q. $y = \{y_1\}$ et par le lemme 9.1.3 $x \in |y|_A = |y_1|_A$. Donc par hypothèse de récurrence il existe une arête de G_{c_1} (donc de G_c) de type A t.q. $x \in e(d)$.

- Si n est un lien porte auxiliaire, alors soit c_1 la prémisse de n . $s(G_{c_1}) < s(G_c)$.

Soit d une arête de type A de G_c telle que $x \in e(d)$. d est aussi une arête de G_{c_1} , et par hypothèse de récurrence, il existe $y_1 \in e(c_1)$ t.q. $x \in |y_1|_A$. Soit alors $y \in e(c)$ t.q. il existe $k \geq 1$ et (si $k \geq 2$) $y_2, \dots, y_k \in e(c_1)$ t.q. $y = y_1 \cup \dots \cup y_k$. Par le lemme 9.1.1 on en déduit que $x \in |y|_A$.

Réciproquement, soit $y \in e(c)$ tel que $x \in |y|_A$. Il existe $k \geq 1$ et $y_1, \dots, y_k \in e(c_1)$ t.q. $y = y_1 \cup \dots \cup y_k$. Par le lemme 9.1.1 on en déduit qu'il existe $i \in \{1, \dots, k\}$ t.q. $x \in |y_i|_A$. Par hypothèse de récurrence il existe alors une arête de G_{c_1} (donc de G_c) de type A t.q. $x \in e(d)$.

- Si n est un lien porte principale, alors soit c_1 la prémisse de n : $s(G_{c_1}) < s(G_c)$.

Soit d une arête de type A de G_c telle que $x \in e(d)$. d est aussi arête de G_{c_1} , et par hypothèse de récurrence il existe $y_1 \in e(c_1)$ t.q. $x \in |y_1|_A$. Soit alors $y \in e(c)$ t.q. $y_1 \in y$. Par le lemme 9.1.1 $x \in |y|_A$.

Réciproquement, soit $y \in e(c)$ tel que $x \in |y|_A$. Il existe $k \geq 1$ et $y_1, \dots, y_k \in e(c_1)$ t.q. $y = \{y_1, \dots, y_k\}$. Par le lemme 9.1.1, il existe $i \in \{1, \dots, k\}$ t.q. $x \in |y_i|_A$. Par hypothèse de récurrence il existe alors une arête d de G_{c_1} (donc de G_c) de type A t.q. $x \in e(d)$. $\square$

Dans la formulation ci-dessus, le lemme précédent est faux si on prend une arête d de type $?A$.

9.1.5. PROPOSITION. Soit R un réseau, c une arête de type C conclusion de R et e une expérience de R . Soit A une occurrence de sous-formule de C t.q. $A \neq ?F$ pour toute formule F , et soient $d_1, \dots, d_k$ les k arêtes ($k \geq 1$) de type A t.q. $\forall i \in \{1, \dots, k\} d_i \in G_c^R$.

Si $e(c) = \{y\}$, alors $|y|_A = e(d_1) \cup \dots \cup e(d_k)$.

Démonstration : On procède par induction sur la cardinalité l de l'ensemble des liens de R traversés par un des chemins issus de l'une des arêtes $d_1, \dots, d_k$ et ayant c comme arête terminale.

Si $l = 0$, alors $C = A$, $k = 1$, $c = d_1$, $|y|_A = \{y\}$ et $e(d_1) = e(c) = \{y\}$.

Autrement, soit n le lien de R dont c est conclusion.

Si $n = \wp$ (resp. $n = \otimes$), alors $C = D \wp E$ (resp. $C = D \otimes E$). Supposons que A soit sous-formule de D , et soit R' un sous-réseau de R ayant l'arête d de type D (prémisse de n dans R) parmi ses conclusions. $\forall i \in \{1, \dots, k\} d_i \in G_d^{R'}$. Si $y = (y_1, y_2)$, alors par le lemme 9.1.3 $|y|_A = |y_1|_A$. La restriction de e à R' est une expérience e' de R' (puisque D est à profondeur 0 dans R). On peut donc appliquer l'hypothèse d'induction à R' , d de type D , e' : puisque $e'(d) = \{y_1\}$, on aura $|y_1|_A = e'(d_1) \cup \dots \cup e'(d_k)$. Le fait que $\forall i \in \{1, \dots, k\} e(d_i) = e'(d_i)$ permet alors de conclure.

Si n est un lien déréluction, alors $C = ?D$, et par hypothèse A est sous-formule de D . Soit R' un sous-réseau de R ayant l'arête d de type D (prémisse de n dans R) parmi ses conclusions. $\forall i \in \{1, \dots, k\} d_i \in G_d^{R'}$. Si $y = \{y'\}$, alors par le lemme 9.1.3 $|y|_A = |y'|_A$. La restriction de e à R' est une expérience e' de R' (puisque D est à profondeur 0 dans R). On peut donc appliquer l'hypothèse d'induction à R' , d de type D , e' : puisque $e'(d) = \{y'\}$, on aura $|y'|_A = e'(d_1) \cup \dots \cup e'(d_k)$. Le fait que $\forall i \in \{1, \dots, k\} e(d_i) = e'(d_i)$ permet alors de conclure.

Si n est un lien affaiblissement, alors $\forall i \in \{1, \dots, k\} e(d_i) = \emptyset$ et $|y|_A = \emptyset$.

Si n est un lien pal (resp. un lien pax), alors $C = !D$ (resp. $C = ?D$). Soit d de type D (resp. de type $?D$) la prémisses de n et soit R' le plus grand sous-réseau de R contenu dans la boîte dont n est une pal (resp. une pax). $\forall i \in \{1, \dots, k\}$ $d_i \in G_d^{R'}$. Soient $e_1, \dots, e_h$ les h expériences de R' t.q. $e_i(d) = \{y_i\}$ et $y = \{y_1, \dots, y_h\}$ (resp. t.q. $e_i(d) = \{y_i\}$ et $y = y_1 \cup \dots \cup y_h$). On peut appliquer l'hypothèse d'induction à R' , d de type D (resp. de type $?D$), $e_i \forall i \in \{1, \dots, h\}$: on aura $|y_i|_A = e_i(d_1) \cup \dots \cup e_i(d_k)$. Par le lemme 9.1.1 (on utilise le fait que dans tous les cas A est sous-formule de D) $|y|_A = |y_1|_A \cup \dots \cup |y_h|_A = e_1(d_1) \cup \dots \cup e_1(d_k) \cup \dots \cup e_h(d_1) \cup \dots \cup e_h(d_k)$. Par le lemme 6.3.2 on a par ailleurs $e(d_i) = e_1(d_i) \cup \dots \cup e_h(d_i) \forall i \in \{1, \dots, h\}$. D'où $|y|_A = e(d_1) \cup \dots \cup e(d_k)$.

Si n est un lien contraction d'arité h , alors $C = ?D$. Soient $c_1, \dots, c_h$ les h prémisses de n de type C , soit R' un sous-réseau de R ayant $c_1, \dots, c_h$ parmi ses conclusions et soit e' la restriction de e à R' . On a $y = y_1 \cup \dots \cup y_h$, avec $e'(c_j) = \{y_j\} \forall j \in \{1, \dots, h\}$. On a $G_c^{R'} = G_{c_1}^{R'} \cup \dots \cup G_{c_h}^{R'} \cup \{n\} \cup \{c\}$. Par conséquent, $\forall i \in \{1, \dots, k\}$ il existe un unique $j \in \{1, \dots, h\}$ t.q. $d_i \in G_{c_j}^{R'}$. On peut donc renommer les arêtes $d_1, \dots, d_k$ de telle sorte que $\{d_1, \dots, d_k\} = \{d_1^1, \dots, d_{l_1}^1, \dots, d_1^h, \dots, d_{l_h}^h\}$ où $d_i^j \in G_{c_j}^{R'}$. On peut appliquer l'hypothèse d'induction à R' , e' , c_j de type $C \forall j \in \{1, \dots, h\}$: on aura $|y_j|_A = e'(d_1^j) \cup \dots \cup e'(d_{l_j}^j)$. Par le lemme 9.1.2 (on utilise le fait que par hypothèse A est sous-formule de D), on a $|y|_A = |y_1|_A \cup \dots \cup |y_h|_A = e'(d_1^1) \cup \dots \cup e'(d_{l_1}^1) \cup \dots \cup e'(d_1^h) \cup \dots \cup e'(d_{l_h}^h) = e'(d_1) \cup \dots \cup e'(d_k) = e(d_1) \cup \dots \cup e(d_k)$. $\square$

9.2 Le cas de la sémantique relationnelle

Les deux propositions suivantes sont vraies aussi pour la sémantique cohérente, mais nous n'en ferons pas usage. L'uniformité de la sémantique cohérente nous fournira en effet un instrument bien plus puissant : la proposition 9.3.5.

9.2.1. PROPOSITION. Soient e et e' deux expériences du réseau R de conclusions respectives γ et γ' .

Si e est une 1-expérience de R et $\gamma = \gamma'$, alors e' est aussi une 1-expérience de R .

Démonstration : Pour prouver que e' est une 1-expérience, il suffira de prouver que e' satisfait la condition (2) de la définition 8.1.1. Ceci n'est qu'une simple conséquence du lemme 9.1.4 : soit a' conclusion de R' et $c' \in G_{a'}^{R'}$ de type $!C$. Si $z' \in e(c')$ n'est pas un singleton, alors $z' \in |e(a')|_{!C} =$

$|e(a)|_C$ (où a est l'arête conclusion de R correspondant à a' , cf. remarque 6.2.2). Donc il existe une arête $c \in G_a^R$ t.q. $z' \in e(c)$, contredisant le fait que e est une 1-expérience de R . Si $e(c) = \emptyset$, on raisonne de la même façon (cf. aussi la démonstration de la proposition 9.2.2). $\square$

9.2.2. PROPOSITION. Soient R et R' deux réseaux de mêmes conclusions. Soit e (resp. e') une expérience de R (resp. de R') de conclusion γ (resp. γ'). Si e' est une n -expérience obsessionnelle simple de R' avec $n > 1$ et si $\gamma = \gamma'$, alors e est une n -expérience obsessionnelle simple de R .

Démonstration : Supposons par absurde que e ne soit pas une n -expérience obsessionnelle de R . Alors (par définition), e ne vérifie pas une des deux conditions de la définition de n -expérience obsessionnelle, c.à.d. qu'on a une des possibilités suivantes :

- (1) il existe une arête a de type atomique X dans R et il existe $x_1, x_2 \in e(a)$ t.q. $x_1 \neq x_2$
- (2) il existe une arête d de type $!D$ de R telle que $e(d) = \emptyset$
- (3) il existe une arête d de type $!D$ de R et $y \in e(d)$ t.q. $\text{card}(y) \neq n$

Remarquons que dans le cas (2), il existe une arête d de type $!D$ de R t.q. $\emptyset \in e(d)$, donc ce cas est en fait un cas particulier du cas (3) avec $y = \emptyset$.

Dans le cas (3), appelons c l'arête de type C conclusion de R telle que $d \in G_c^R$ et soit $e(c) = \{t\}$. Par le lemme 9.1.4 $y \in |t|_D$. Si on appelle c' l'arête conclusion de R' de type C correspondant à c , on a $e(c) = e'(c')$. En appliquant de nouveau le lemme 9.1.4 on déduit l'existence de d' arête de R' de type $!D$ t.q. $y \in e'(d')$, et e' n'est pas une n -expérience obsessionnelle.

Dans le cas (1), appelons c l'arête de type C conclusion de R et telle que $a \in G_c^R$ et soit $e(c) = \{t\}$. Par le lemme 9.1.4, $x_1 \in |t|_X$ et $x_2 \in |t|_X$. Si on appelle c' l'arête de R' de type C correspondant à c , on a $e(c) = e'(c')$. En appliquant de nouveau le lemme 9.1.4 on déduit l'existence de a'_1 arête de R' de type atomique X t.q. $x_1 \in |t|_X$ et l'existence de a'_2 arête de R' de type atomique X t.q. $x_2 \in |t|_X$. Remarquons que à priori rien ne garantit que $a'_1 = a'_2$. Mais alors même si e' est une n -expérience obsessionnelle, ce n'est certainement pas une n -expérience obsessionnelle simple. $\square$

La conjecture qui suit (et dont nous n'avons pas de preuve pour l'instant) est, dans le cadre de la sémantique cohérente, une conséquence de la propriété d'uniformité (cf. proposition 9.3.5).

9.2.3 Conjecture

Soient R et R' deux réseaux de mêmes conclusions. Soit e (resp. e') une expérience de R (resp. de R') de conclusion γ (resp. γ').

Si e' est une n -expérience obsessionnelle (quelconque) de R' avec $n > 1$ et si $\gamma = \gamma'$, alors il existe une n -expérience obsessionnelle e_1 de R de conclusion γ .

9.3 Le cas de la sémantique cohérente

Nous avons déjà dit que la sémantique cohérente est uniforme, et que c'est précisément ce qui la différencie de la sémantique relationnelle.

Le lemme suivant est une façon d'exprimer cette uniformité, qui dans notre contexte prendra surtout la forme de la proposition 9.3.5.

9.3.1. LEMME. (Propriété d'uniformité) Supposons que le réseau R soit une boîte B de conclusions l'arête d de type $!D$ et les arêtes $d_1, \dots, d_k$ de types respectifs $?D_1, \dots, ?D_k$. Appelons d' (resp. $d'_1, \dots, d'_k$) la prémisse de la porte principale de B (resp. les prémisses des pax de conclusions $d_1, \dots, d_k$ de B) : $d', d'_1, \dots, d'_k$ sont les conclusions de R_B . Soit e une expérience de R t.q. $e(d) = \{\{y_1, \dots, y_n\}\}$, et soient $e_1, \dots, e_n$ les n expériences de R_B t.q. $\forall i \in \{1, \dots, n\} \ e_i(d') = \{y_i\}$. Appelons enfin $y_i, t_1^{i_1}, \dots, t_k^{i_2}$ la conclusion de l'expérience e_i de $R_B \ \forall i \in \{1, \dots, n\}$.

Si, pour $i_1, i_2 \in \{1, \dots, n\}$ on a $y_{i_1} = y_{i_2}$, alors $\forall j \in \{1, \dots, k\}$ on a $t_j^{i_1} = t_j^{i_2}$.

Démonstration : Les conclusions des expériences $e_1, \dots, e_n$ de R_B sont éléments d'une même clique de l'espace cohérent $\mathcal{A} = \mathcal{D} \wp ?\mathcal{D}_1 \wp \dots \wp ?\mathcal{D}_k$, donc $(y_{i_1}, t_1^{i_1}, \dots, t_k^{i_1}) \subset (y_{i_2}, t_1^{i_2}, \dots, t_k^{i_2})(\mathcal{A})$. Par ailleurs, $\forall j \in \{1, \dots, k\}$ l'étiquette associée par e à la conclusion d_j de type $?D_j$ de R est $t_j^{i_1} \cup \dots \cup t_j^{i_n}$ qui est un élément de $|\mathcal{D}_j| = \mathcal{D}_j^\perp$. Donc en particulier $t_j^{i_1} \cup t_j^{i_2}$ est une clique de $\mathcal{D}_j^\perp$, c.à.d. que $t_j^{i_1} \smile t_j^{i_2} (?D_j)$. Mais alors, pour aucun $j \in \{1, \dots, k\}$ on aura $t_j^{i_1} \wedge t_j^{i_2} (?D_j)$. Puisque par ailleurs $y_{i_1} = y_{i_2}$ on n'a pas non plus $y_{i_1} \wedge y_{i_2}$. La seule possibilité pour que $(y_{i_1}, t_1^{i_1}, \dots, t_k^{i_1}) \subset (y_{i_2}, t_1^{i_2}, \dots, t_k^{i_2})(\mathcal{A})$ est alors que $\forall j \in \{1, \dots, k\} \ t_j^{i_1} = t_j^{i_2}$. $\square$

9.3.2. LEMME. Soit n un entier strictement positif. Soit e' une expérience du réseau R de résultat γ t.q. pour toute arête c de type $!C$ de R , si on appelle c' (de type C) la prémisse du lien pal dont c est conclusion, on a :

(1) $e'(c) \neq \emptyset$

(2) pour tout $y \in e'(c)$, il existe $z \in e'(c')$ t.q. $y = \{n[z]\}$.

Alors il existe une n -expérience obsessionnelle e de R de conclusion γ .

Démonstration : On procède par induction sur une séquentialisation du réseau R .

Si R est un axiome, le lemme est vrai.

Autrement, soit n un lien terminal de R , c.à.d. tel qu'il existe une séquentialisation de R dont la dernière règle correspond à n . Si R n'est pas une boîte, alors $n = \otimes, \wp, ?de, ?co, ?w$, et le résultat est une application de l'hypothèse d'induction. En guise d'exemple, nous allons considérer les cas $n = \otimes$ et $n = ?co$. Si $n = \otimes$, appelons R_1 et R_2 les deux sous-réseaux de R de conclusions respectives Γ_1, A_1 et Γ_2, A_2 obtenus à partir de R en effaçant le lien $\otimes$ terminal n et sa conclusion. Appelons e'_1 (resp. e'_2) la restriction de e' au réseau R_1 (resp. R_2). Par hypothèse d'induction, il existe une n -expérience obsessionnelle e_1 (resp. e_2) de R_1 (resp. R_2) de même conclusion que e'_1 (resp. e'_2). La n -expérience obsessionnelle e de R cherchée est alors obtenue de façon évidente à partir de e_1 et e_2 .

Si $n = ?co$, appelons R_1 le sous-réseau de R obtenu à partir de R en effaçant le lien $?co$ terminal n et sa conclusion. Appelons e'_1 la restriction de e' au réseau R_1 . Par hypothèse d'induction, il existe une n -expérience obsessionnelle e_1 de R_1 de même conclusion que e'_1 . L'expérience e de R cherchée est alors obtenue de façon évidente à partir de e_1 . Il est important de remarquer, dans ce cas, que la compatibilité des étiquettes associées par e_1 aux arêtes conclusions de R_1 et prémisses de n est assurée par l'hypothèse d'induction : ce sont les mêmes étiquettes que l'expérience e' associe à ces arêtes.

Considérons enfin le cas où R est une boîte B de conclusions l'arête d de type $!D$ et les arêtes $d_1, \dots, d_k$ de types respectifs $?D_1, \dots, ?D_k$. Appelons d' (resp. $d'_1, \dots, d'_k$) la prémisses de la porte principale de B (resp. les prémisses des pax de conclusions $d_1, \dots, d_k$ de B) : $d', d'_1, \dots, d'_k$ sont les conclusions de R_B . Par hypothèse $e'(d) = \{\{n[y]\}\}$. Il existe donc n expériences $e'_1, \dots, e'_n$ de R_B t.q. $\forall i \in \{1, \dots, n\} e'_i(d') = \{y\}$. Soit alors $y, t_1^i, \dots, t_k^i$ la conclusion de l'expérience e'_i de R_B $\forall i \in \{1, \dots, n\}$. Nous allons prouver que :

(a) $\forall i_1, i_2 \in \{1, \dots, n\} \forall j \in \{1, \dots, k\}$ on a $t_j^{i_1} = t_j^{i_2}$

(b) $\forall i \in \{1, \dots, n\}$ il existe une n -expérience obsessionnelle e_i de R_B de même conclusion que e'_i .

(a) est une application de la propriété d'uniformité (lemme 9.3.1). Par hypothèse d'induction, il suffira de prouver que $\forall i \in \{1, \dots, n\}$ l'expérience e'_i de R_B satisfait les hypothèses (1) et (2) du lemme. Si pour une certaine arête c de R_B $e'_i(c) = \emptyset$, alors il existe forcément une arête g de type $!G$ de

R_B telle que $\emptyset \in e'_i(g)$. Mais par le lemme 6.3.2, on aurait alors $\emptyset \in e'(g)$, ce qui contredit l'hypothèse (2) du lemme. De même, supposons qu'il existe une arête c de type $!C$ et $y \in e'_i(c)$ t.q. $\forall z \in e'_i(c') \ y \neq \{n[z]\}$, où c' est l'arête de R_B prémisses du lien promotion dont c est conclusion. Dans ce cas, la cardinalité de y (en tant que multiensemble) est différente de n ou bien il existe $z_1, z_2 \in e'_i(c')$ t.q. $z_1 \neq z_2$ et $z_1, z_2 \in y$. Par le lemme 6.3.2, dans les deux cas il existe $y \in e'(c)$ t.q. $\forall z \in e'(c') \ y \neq \{n[z]\}$, ce qui contredit l'hypothèse du lemme.

Soit donc e_1 une n -expérience obsessionnelle de R_B de même résultat que e'_1 (et donc que toutes les e'_i par la propriété (a) ci-dessus). On définit l'expérience e de R comme l'unique n -expérience obsessionnelle de R obtenue à partir de la n -expérience obsessionnelle e_1 de R_B . $\square$

9.3.3. REMARQUE. (i) Si le théorème de Duquesne et Van de Wiele (cf. remarque 6.2.7) s'étend à *MELL*, alors le lemme précédent permet d'affirmer (à la lumière de la proposition 8.1.6) que dans le cas cohérent on peut substituer les conditions (1) et (2) de la définition de n -expérience obsessionnelle par l'hypothèse du lemme.

(ii) Soit e (resp. e') une n -expérience obsessionnelle du réseau R (resp. du réseau R') et soit a (resp. a') une arête de R (resp. de R') conclusion du lien m (resp. m') dont les prémisses sont $a_1, \dots, a_k$ (resp. $a'_1, \dots, a'_k$), avec $k \geq 0$.

Si m et m' sont des liens du même type et si $\forall i \in \{1, \dots, k\} \ e(a_i) = e'(a'_i)$, alors $e(a) = e'(a')$.

Remarquons au passage que ceci est faux si e et e' ne sont pas n -obsessionnelles.

Démonstration : On montre (ii). Soit $x \in e(a)$ (resp. $x' \in e'(a')$). $\forall i \in \{1, \dots, k\}$ il existe $x_i \in e(a_i)$ (resp. $x'_i \in e'(a'_i)$) t.q. x (resp. x') s'obtient selon les modalités relatives au type du lien m (resp. m') à partir de $(x_1, \dots, x_k)$ (resp. de $(x'_1, \dots, x'_k)$). Par la proposition 8.1.6 $\forall i \in \{1, \dots, k\}$, de $e(a_i) = e'(a'_i)$ on peut déduire $x_i = x'_i$. Mais alors puisque m et m' sont de même type on aura forcément $x = x'$. Par ailleurs, puisque $\forall i \in \{1, \dots, k\} \ e(a_i) = e'(a'_i)$, la proposition 8.1.6 permet d'affirmer que a_i et a'_i ont la même profondeur. Donc a et a' ont aussi la même profondeur. Et alors $e(a) = e'(a')$.

La propriété est fausse si on omet l'hypothèse e et e' n -expérience obsessionnelles : prenons par exemple $k = 2$, m et m' de type $\otimes$, $e(a_1) = e'(a'_1) = \{x_1, x_2\}$ et $e(a_2) = e'(a'_2) = \{y_1, y_2\}$, avec $x_1 \neq x_2$ et $y_1 \neq y_2$. On peut très bien avoir $e(a) = \{(x_1, y_1), (x_2, y_2)\}$ et $e'(a') = \{(x_1, y_2), (x_2, y_1)\}$. $\square$

9.3.4. PROPOSITION. Soit R un réseau, $\gamma \in [R]$ et n un entier strictement positif. Les deux affirmations suivantes sont équivalentes :

- (i) Pour toute arête a de type A conclusion de R et pour toute (occurrence de) sous-formule $!C$ de A , si x est l'étiquette de γ associée à a , on a :
si $z \in |x|_{!C}$, alors il existe $t \in |\mathcal{C}|$ t.q. $z = \{n[t]\}$.
- (ii) Il existe une n -expérience obsessionnelle e_1 de R de résultat γ .

Démonstration : Soit e_1 une n -expérience obsessionnelle de R de résultat γ , et a une arête de type A conclusion de R . Soit $!C$ une (occurrence de) sous-formule de A et x l'étiquette que e_1 associe à a . Par la proposition 9.1.5, $|x|_{!C} = e_1(c_1) \cup \dots \cup e_1(c_k)$, où $c_1, \dots, c_k$ sont les k arêtes ($k \geq 1$) de type $!C$ t.q. $\forall i \in \{1, \dots, k\}$ on a $c_i \in G_a^R$. Si $z \in |x|_{!C}$, alors il existe $i \in \{1, \dots, k\}$ t.q. $z \in e_1(c_i)$. Par la remarque 8.1.7 (i), on a donc bien $z = \{n[t]\}$, pour un certain $t \in |\mathcal{C}|$.

Réciproquement, soit e une expérience de R de résultat γ . On montre que e vérifie les hypothèses du lemme 9.3.2, en appliquant le lemme 9.1.4.

Supposons par absurde que e ne satisfasse pas les hypothèses du lemme 9.3.2. Il existe alors une arête c de type $!C$ conclusion d'un lien porte principale de prémisses c' satisfaisant une des deux conditions suivantes :

- (1) $e(c) = \emptyset$
- (2) $\exists y \in e(c)$ t.q. $\forall z \in e(c') y \neq \{n[z]\}$.

Dans le cas (1), e n'est pas définie sur l'arête c de R . Par le lemme 6.3.3, il existe donc une boîte B de R t.q. si on appelle d la conclusion de type $!D$ de la pal de B , alors $e(d) = \{m[\emptyset]\}$, pour un certain entier non nul m . Appelons a de type A la conclusion de R telle que $d \in G_a^R$ et soit $x \in |\mathcal{A}|$ tel que $e(a) = \{x\}$. Le lemme 9.1.4 appliqué à e donne $\emptyset \in |x|_{!D}$, ce qui contredit (i).

Dans le cas (2), on a deux possibilités :

- (2.1) $\exists y \in e(c)$ t.q. $\text{card}(y) \neq n$
- (2.2) $\exists y \in e(c)$ t.q. $\text{card}(y) = n$ et $\exists z_1, z_2 \in e(c'), z_1 \neq z_2$ t.q. $z_1, z_2 \in y$.

Dans les deux cas, appelons a de type A la conclusion de R telle que $c \in G_a^R$ et soit $x \in |\mathcal{A}|$ tel que $e(a) = \{x\}$. Le lemme 9.1.4 appliqué à e donne $y \in |x|_{!C}$, ce qui contredit à nouveau (i). $\square$

9.3.5. PROPOSITION. Soient R et R' deux réseaux de mêmes conclusions. Soit e (resp. e') une expérience de R (resp. de R') de conclusion γ (resp. γ'). Si e' est une n -expérience obsessionnelle de R' et si $\gamma = \gamma'$, alors il existe une expérience e_1 de R de conclusion $\gamma = \gamma'$ qui est une n -expérience obsessionnelle.

Démonstration : Il suffit de montrer que le résultat γ de e satisfait (i) de la proposition 9.3.4. Ceci est à peu près évident, puisque e' est n -obsessionnelle et son résultat $\gamma' = \gamma$ satisfait donc (i) de la proposition 9.3.4. $\square$

9.3.6. REMARQUE. Il est probable qu'au prix d'un (petit ?) effort technique on puisse prouver que dans les énoncés des deux propositions précédentes, e elle-même est une n -expérience obsessionnelle. D'ailleurs, si le théorème de Duquesne et Van de Wiele (cf. remarque 6.2.7) peut s'étendre à $MELL$, alors ceci est évident (puisque $e = e_1$).

Chapitre 10

Bilan sur l'injectivité : résultats et conjectures

Arrivés à ce stade de l'analyse, nous faisons un bilan. Nous montrons que tant pour la sémantique cohérente que pour la sémantique relationnelle, l'interprétation d'un réseau standard R (sans affaiblissements) permet de reconstruire R “aux liens axiome près” (théorèmes 10.1.1 et 10.2.3).

Pour la sémantique relationnelle, nous montrons qu'un réseau standard sans affaiblissements est seul dans sa “classe d'équivalence sémantique” lorsqu'il ne contient pas de contractions (théorème 10.1.4), de même que lorsqu'il ne contient pas de boîtes (théorème 10.1.3).

Pour la sémantique cohérente, nous prouvons qu'un réseau R standard et sans affaiblissements est seul dans sa “classe d'équivalence sémantique” lorsqu'il existe une n -expérience obsessionnelle injective de R (théorème 10.2.4). C'est précisément la question de l'existence d'une telle expérience qui nous occupera dans les chapitres 11, 12, 13 et 14.

Avant de faire le point, nous allons prouver la proposition suivante, qui utilise une caractérisation très simple des boîtes exponentielles en l'absence d'affaiblissements.

10.0.7. PROPOSITION. Soient R et R' deux réseaux de mêmes conclusions ne contenant aucun lien affaiblissement.

Si $SPL(R) = SPL(R')$, alors $R = R'$.

Démonstration : Nous utiliserons la caractérisation suivantes des boîtes :

Soit l (resp. m) un lien pax ou pal de profondeur p dans le réseau T (T ne contenant aucun lien affaiblissement). Les liens l et m sont deux portes d'une même boîte si et seulement si il existe un chemin orienté quelconque (pas forcément comme dans la définition 2.1.1) ayant l comme premier lien et m comme dernier lien et dont *tous* les autres liens sont à profondeur strictement supérieure à p .

Nous laissons au lecteur le soin de se convaincre que la propriété ci-dessus est effectivement caractéristique de deux portes quelconques d'une même boîte.

La conclusion est alors une simple conséquence du fait que les chemins quelconques de $SPL(T)$ sont exactement ceux de T , pour tout réseau T . $\square$

10.1 Bilan pour la sémantique relationnelle

Relativement à la sémantique relationnelle, nous avons le

10.1.1. THÉORÈME. Soient R et R' deux réseaux de mêmes conclusions. Si $[R] = [R']$, alors $SPLO(R) = SPLO(R')$ et $SPP(R) = SPP(R')$.

Démonstration : Soit h l'arité maximale des liens contraction de R et R' , et soit $n > h$ si $h \geq 2$ et $n = 1$ autrement. Il existe certainement une n -expérience obsessionnelle simple de R (nous sommes dans le cas relationnel), puisque R est sans coupures (cf. (iv) de la remarque 6.2.2). Soit donc e_n une telle expérience. Puisque $[R] = [R']$, il existe une expérience e'_n de R' de même conclusion que e_n . Si $n = 1$ (resp. $n > 1$), alors la proposition 9.2.1 (resp. 9.2.2) permet d'affirmer que e'_n est une n -expérience obsessionnelle simple de R' . Le corollaire 8.3.9 permet alors de conclure que $SPLO(R) = SPLO(R')$.

Montrons maintenant que $SPP(R) = SPP(R')$. Soit e_1 une 1-expérience injective de R (qui existe certainement pour les raisons évoquées ci-dessus). Puisque $[R] = [R']$, il existe une expérience e'_1 de R' de même conclusion que e_1 . La proposition 9.2.1 permet d'affirmer que e'_1 est une 1-expérience de même conclusion que e . On en déduit (grâce au corollaire 8.2.3) que $SPPO(R) = SPPO(R')$ et $e_1|_{SPPO(R)} = e'_1|_{SPPO(R')}$. L'injectivité de e_1 (donc de e'_1) permet de conclure que $SPP(R) = SPP(R')$. $\square$

10.1.2. REMARQUE. Il existe des réseaux R et R' pour lesquels $SPP(R) = SPP(R')$, $SPLO(R) = SPLO(R')$, et pourtant $R \neq R'$. Ceci est essentiellement dû au fait que G_a^R et GP_a^R sont des *classes d'équivalences* d'arbres,

comme nous l'avons précisé dans la remarque 7.2.4 (cf. aussi la preuve du théorème 10.1.4).

10.1.3. THÉORÈME. Soient R et R' deux réseaux de mêmes conclusions, et supposons que dans les types de ces conclusions il n'y ait aucune occurrence de l'exponentielle bien sûr.

Si $[R] = [R']$, alors $R = R'$.

Démonstration : On applique le théorème précédent. Dans notre cas, puisque R ne contient pas de boîtes on aura : $R = SPP(R) = SPP(R') = R'$. $\square$

10.1.4. THÉORÈME. Soient R et R' deux réseaux de mêmes conclusions, et supposons que R ne contienne aucun lien affaiblissement ni aucun lien contraction.

Si $[R] = [R']$, alors $R = R'$.

Démonstration : Puisqu'il n'y a pas de contractions dans R (et donc dans R' puisque par le théorème 10.1.1 $SPP(R) = SPP(R')$), les graphes des arêtes conclusion de R et de R' ne sont pas des classes d'équivalence d'arbres mais tout simplement des arbres.

Reprenons à la lumière de cette remarque la démonstration du théorème 10.1.1. Soient a et a' deux arêtes conclusions de R et de R' respectivement, et qui se correspondent grâce à l'égalité $[R] = [R']$ (cf. remarque 6.2.2). On aura par le théorème 10.1.1 $G_a^R = G_{a'}^{R'}$. Soit maintenant e_1 une 1-expérience injective de R et e'_1 la 1-expérience injective de R' de même conclusion que e_1 dont il est question dans la preuve du théorème 10.1.1. On aura en particulier $e_1(a) = e'_1(a')$. *Puisqu'il n'y a pas de liens contraction dans R et dans R'* l'égalité $G_a^R = G_{a'}^{R'}$ est une égalité d'arbres, et l'égalité $e_1|_{GP_a^R} = e'_1|_{GP_{a'}^{R'}}$ induit une *unique* bijection entre les arêtes de type atomique de GP_a^R et celles de $GP_{a'}^{R'}$, c.à.d. entre les arêtes de type atomique de G_a^R et celles de $G_{a'}^{R'}$.

Remarquons que lorsque ces graphes sont des classes d'équivalence il y a en général plusieurs bijections possibles.

Cette unique bijection entre les arêtes de $SPLO(R)$ et celles de $SPLO(R')$ permet d'affirmer que $SPL(R) = SPL(R')$, et donc par la proposition 10.0.7 que $R = R'$. $\square$

10.1.5. REMARQUE. Bien que les deux théorèmes précédents soient des résultats partiels, ils le sont peut-être moins qu'ils n'en ont l'air : R et R' sont en effet standard, et il se pourrait que ce soient les formes normales de réseaux qui, eux, contiennent des boîtes ou des liens $?co$ et $?w$.

Si la conjecture 9.2.3 est vraie, on en déduit la

10.1.6 Conjecture

Soient R et R' deux réseaux de mêmes conclusions ne contenant aucun lien affaiblissement. Si $[R] = [R']$, alors $R = R'$.

“Démonstration :” On suppose qu'on a prouvé la conjecture 9.2.3.

Soit h l'arité maximale des liens contraction de R et R' , et soit $n > h$ (on suppose que $h \geq 2$, puisque dans le cas contraire le théorème 10.1.1 résout la question).

Soit alors e_n une n -expérience obsessionnelle injective (qui existe pour les raisons usuelles) de R . Puisque $[R] = [R']$, il existe une expérience e' de R' de même conclusion que e_n . La conjecture 9.2.3 nous dit qu'il existe alors e'_n de même conclusion que e_n et e' qui est n -obsessionnelle. Le corollaire 8.3.9 permet de conclure que $SPLO(R) = SPLO(R')$ et $e_n|_{SPLO(R)} = e'_n|_{SPLO(R')}$. De l'injectivité de e_n (donc de e'_n) on déduit $SPL(R) = SPL(R')$, et la proposition 10.0.7 donne enfin $R = R'$. $\square$

10.2 Bilan pour la sémantique cohérente

Contrairement au cas relationnel, l'existence d'une expérience simple ou injective n'est pas évidente dans le cas cohérent, et ce à cause de “l'uniformité” de la sémantique cohérente : même pour les expériences obsessionnelles, même pour les réseaux standard, nous ne sommes pas assurés à priori que les étiquettes des prémisses d'un lien contraction satisfassent la contrainte requise par la définition 6.2.1. On peut voir cette difficulté comme le “prix à payer” pour la simplicité avec laquelle nous avons prouvé la proposition 9.3.5.

On va donc commencer par démontrer qu'il existe pour tout réseau de *MELL* une expérience simple cohérente (proposition 10.2.2).

Notation : Si e est une n -expérience obsessionnelle du réseau R et a est une arête de R , on notera à partir de maintenant $|e(a)|$ l'unique élément de $e(a)$. Nous n'avons pas introduit avant cette notation (qui devient désormais

indispensable) pour ne pas créer de confusion avec la notion de projection de la définition 6.3.5.

10.2.1. LEMME. Soit R un réseau, a et a' deux arêtes de R de type A . Soit e une n -expérience obsessionnelle de R .

Si $\forall \alpha \in G_a$, et $\forall \alpha' \in G_{a'}$, avec α et α' de même type atomique, on a $|e(\alpha)| = |e(\alpha')|$, alors $|e(a)| \sim |e(a')|(\mathcal{A})$.

Démonstration : Par induction sur $p = s(G_a) + s(G_{a'})$, où pour toute arête b de R $s(G_b)$ = nombre d'arêtes de G_b .

Si $p = 2$, alors a et a' sont conclusions de deux liens axiome et par hypothèse $|e(a)| \sim |e(a')|(\mathcal{A})$.

Autrement il va falloir distinguer différents cas pour la formule A :

Si $A = B \otimes C$ (resp. $A = B \wp C$), appelons b et b' les prémisses de type B des liens $\otimes$ (resp. $\wp$) dont a et a' sont conclusions, et c et c' les prémisses de type C de ces mêmes liens. L'hypothèse d'induction appliquée à G_b , $G_{b'}$ et à G_c , $G_{c'}$ donne $|e(b)| \sim |e(b')|(\mathcal{B})$ et $|e(c)| \sim |e(c')|(\mathcal{C})$. Tant dans le cas du $\otimes$ que dans le cas du $\wp$, ceci permet de conclure $|e(a)| \sim |e(a')|(\mathcal{A})$.

Si $A = !B$, appelons b et b' les prémisses de type B des liens promotion dont a et a' sont conclusions. L'hypothèse d'induction appliquée à G_b , $G_{b'}$ donne $|e(b)| \sim |e(b')|(\mathcal{B})$. Si $|e(b)| = |e(b')|$, alors $|e(a)| = |e(a')|$. Si par contre $|e(b)| \prec |e(b')|(\mathcal{B})$, alors tout multienemble contenant $|e(b)|$ et $|e(b')|$ n'est pas une clique de $\mathcal{B}$: ce sera le cas notamment pour $|e(a)| \cup |e(a')|$, et donc $|e(a)| \prec |e(a')|(\mathcal{A})$. Dans tous les cas on aura donc bien $|e(a)| \sim |e(a')|(\mathcal{A})$.

Si $A = ?B$, lorsque a ou/et a' est conclusion d'un lien affaiblissement on a $\emptyset \sim |e(a')|(?B)$ ou/et $|e(a)| \sim \emptyset(?B)$. Lorsque ni a ni a' n'est conclusion d'un affaiblissement, appelons $b_1, \dots, b_k$ les prémisses des k liens déréluction au-dessus de a (on aura éventuellement $k = 1$), et $b'_1, \dots, b'_h$ les prémisses des h liens déréluction au-dessus de a' (on aura éventuellement $h = 1$). Remarquons que tous les b_i peuvent être parmi les b'_j ou/et réciproquement. Puisque $\forall i \in \{1, \dots, k\}$ et $\forall j \in \{1, \dots, h\}$ les arêtes b_i et b'_j sont de type B , l'hypothèse d'induction appliquée à G_{b_i} et $G_{b'_j}$ donne $|e(b_i)| \sim |e(b'_j)|(\mathcal{B})$. Nous savons que $\forall x \in |e(a)| \exists i \in \{1, \dots, k\}$ t.q. $x = |e(b_i)|$ et $\forall y \in |e(a')| \exists j \in \{1, \dots, h\}$ t.q. $y = |e(b'_j)|$. On déduit alors de ce qui précède que $\forall x \in |e(a)|$ et $\forall y \in |e(a')|$ $x \sim y(\mathcal{B})$ i.e. $x \subset y(\mathcal{B}^\perp)$. Donc $|e(a)| \cup |e(a')|$ est une clique de $\mathcal{B}^\perp$ c.à.d. $|e(a)| \sim |e(a')|(\mathcal{B})$. $\square$

10.2.2. PROPOSITION. Soit R un réseau. $\forall n \geq 1$ il existe une expérience obsessionnelle simple de degré n de R .

Démonstration : Soit $\mathcal{X}$ un espace cohérent et $x \in |X|$. Si on interprète toutes les formules atomique de R par l'espace $\mathcal{X}$ et si on associe à tous les axiomes l'étiquette x , alors le lemme précédent montre bien qu'on peut toujours effectuer une contraction entre deux arêtes différentes et de même type.

Plus formellement, on peut procéder par induction sur une séquentialisation de R , le seul cas significatif étant celui de la contraction terminale que l'on règle en utilisant le lemme précédent. $\square$

On va maintenant faire le point sur la sémantique cohérente.

10.2.3. THÉORÈME. Soient R et R' deux réseaux de mêmes conclusions. Si $[R] = [R']$, alors $SPLO(R) = SPLO(R')$.

Démonstration : Presque identique à celle du théorème 10.1.1. Il faut simplement invoquer la proposition 10.2.2 pour l'existence d'une expérience simple de R et la proposition 9.3.5 pour la préservation de l'obsessionnalité. $\square$

Reste la question de savoir si pour la sémantique cohérente la conjecture 7.1.3 est vraie ou fausse. Nous allons attaquer la question dans le chapitre suivant en commençant par essayer de montrer l'existence d'une expérience injective cohérente. Il en découlerait par l'argument désormais usuel que si $[R] = [R']$, alors $SPL(R) = SPL(R')$, ce qui en l'absence d'affaiblissements permettrait de conclure que $R = R'$ (grâce à la proposition 10.0.7). En d'autres termes, nous venons de dire que si il existe, pour le réseau R , une expérience n -obsessionnelle (avec n suffisamment grand) et injective, alors R est "seul" dans sa classe d'équivalence sémantique : c'est ce que dit le théorème qui suit, où $[R]_{coh}$ est l'interprétation cohérente multiensembliste de R .

10.2.4. THÉORÈME. Soit R un réseau, et h l'arité maximale de ses liens contraction. Soit $n > \sup(1, h)$. Soit R' un réseau de même conclusion que R .

Si il existe une n -expérience obsessionnelle et injective de R , alors lorsque $[R]_{coh} = [R']_{coh}$ on a forcément $SPL(R) = SPL(R')$. Si, en outre, R ne contient pas de liens affaiblissement, alors $R = R'$.

Démonstration : Soit e_n la n -expérience obsessionnelle et injective de R de conclusion γ dont on suppose l'existence. Puisque $[R]_{coh} = [R']_{coh}$, par la proposition 9.3.5 il existe une n -expérience obsessionnelle e'_n de R'

de conclusion γ . Par le corollaire 8.3.9, on aura $SPLO(R) = SPLO(R')$ et $e_n|_{SPLO(R)} = e'_n|_{SPLO(R')}$. On déduira alors de l'injectivité de e_n que $SPL(R) = SPL(R')$. Lorsque R ne contient pas de liens $?w$, la proposition 10.0.7 permet de conclure que $R = R'$. $\square$

10.2.5. REMARQUE. Il sera beaucoup question, dans la suite, de l'existence ou non d'une expérience n -obsessionnelle injective cohérente pour un réseau R donné d'un certain fragment F de LL (sans affaiblissements).

Il est intéressant de remarquer que par la remarque 7.1.4, une réponse positive à cette question aurait *aussi* comme conséquence immédiate l'injectivité de la sémantique relationnelle pour F .

Chapitre 11

Elimination des boîtes

Le théorème 10.2.4 affirme que pour montrer qu'un réseau R (sans affaiblissements) est "seul dans sa classe d'équivalence sémantique", il suffit de trouver pour R une expérience obsessionnelle cohérente et injective. Nous allons dorénavant nous concentrer sur cette question, et nous ne parlerons presque plus de sémantique relationnelle, restant entendu que tout résultat d'injectivité établi pour la sémantique cohérente multiensembliste restera vrai pour la sémantique relationnelle (par les remarques 7.1.4 et 10.2.5).

Soit R un réseau. Si à tout lien axiome l de conclusion α_l et $\alpha_l^\perp$ on associe un élément x_l de la trame de l'espace cohérent $\mathcal{A}$ de telle sorte que si $l \neq l'$ alors $x_l \neq x_{l'}$, qu'est-ce qui pourrait s'opposer à l'existence d'une n -expérience obsessionnelle e_n de R t.q. pour tout lien axiome l on ait $|e_n(\alpha_l)| = |e_n(\alpha_l^\perp)| = x_l$?

La définition 6.2.1 d'expérience montre bien qu'en l'absence de coupures, les seules contraintes qu'une telle expérience doit satisfaire portent sur les étiquettes des prémisses et des conclusions des liens *pax*, *pal* et *?co*.

Nous montrons, dans ce chapitre, que la présence des boîtes n'a rien à voir avec la question que nous nous posons (comme le suggérait la remarque 8.1.8) : si il existe une expérience (obsessionnelle) injective pour le "linéarisé" $L(R)$ (qui ne contient plus de boîtes) du réseau R , alors il existe une expérience obsessionnelle injective de R (proposition 11.2.4).

11.1 Réduction au cas des 1-expériences

On va d'abord montrer que si il existe une 1-expérience injective d'un réseau R , alors il existe aussi une n -expérience obsessionnelle injective pour R : celle qui est "induite" par la 1-expérience de départ (proposition 11.1.6).

11.1.1. DÉFINITION. Soit R un réseau et e_1 une 1-expérience injective de R . On appelle n -expérience obsessionnelle induite par e_1 toute n -expérience obsessionnelle e_n de R telle que : pour toute arête atomique a de R $|e_n(a)| = e_1(a)$.

11.1.2. REMARQUE. (i) L'existence d'une n -expérience obsessionnelle induite par une 1-expérience injective donnée n'est pas à priori évidente. Par contre, si elle existe il est bien clair qu'elle est unique (par le lemme 8.1.4), et qu'elle est injective.

(ii) Que le lecteur qui pense que si e_n est la n -expérience obsessionnelle induite par la 1-expérience injective e_1 de R alors pour toute arête a , $|e_n(a)| = e_1(a)$, se détrompe : pour ce faire, il n'aura qu'à supposer que $|e_n(c')| = e_1(c')$, où c' est la prémisse d'un lien bien sûr de R de conclusion c , et remarquer que $|e_n(c)| = \{n[|e_n(c')|]\}$, tandis que $e_1(c) = \{e_1(c')\}$.

11.1.3. LEMME. Soient a et a' deux arêtes différentes de même type A du réseau R , et soit $n > \sup(1, h)$ où h est l'arité maximale des liens contraction de R . Soit e_1 une 1-expérience injective de R et e_n la n -expérience obsessionnelle de R induite par e_1 . On a :

- (1) Si il existe une arête α de type atomique t.q. $\alpha \in G_a$ ou bien $\alpha \in G_{a'}$, alors $|e_n(a)| \neq |e_n(a')|$
- (2) Autrement, soit $G_a = G_{a'}$ et alors $|e_n(a)| = |e_n(a')|$, soit $G_a \neq G_{a'}$ et alors $|e_n(a)| \prec |e_n(a')|$.

Démonstration : (1) Soit $x = |e_n(a)|$ (resp. $x' = |e_n(a')|$), et supposons par exemple que $\alpha \in G_a$ soit de type X .

Si $a' \in G_a$ ou bien si $a \in G_{a'}$, alors forcément $A = B$ (n'oublions pas que a et a' sont deux arêtes de même type), et puisque $n > 1$ $x \neq x'$.

Autrement on aura $G_a \cap G_{a'} = \emptyset$. Donc $\alpha \notin G_{a'}$. Pour toute arête $\beta \in G_{a'}$ atomique $|e_n(\beta)| \neq |e_n(\alpha)|$ (par définition d'expérience injective). Par le lemme 9.1.4 on aura alors $|e_n(\alpha)| \in |x|_X$ et $|e_n(\alpha)| \notin |x'|_X$. D'où $x \neq x'$.

(2) Si une telle arête α n'existe pas, alors aucune arête de G_a ou de $G_{a'}$ n'est une arête de type atomique. Il y a donc dans G_a comme dans $G_{a'}$ des liens affaiblissement dont les conclusions sont les seules "feuilles" de G_a et $G_{a'}$.

Si $G_a = G_{a'}$, alors (puisque e_n est n -obsessionnelle) $e_n(a) = e_n(a')$ (et en particulier $|e_n(a)| = |e_n(a')|$).

Si $G_a \neq G_{a'}$, on va procéder par induction sur $k = s(G_a) + s(G_{a'})$, la somme du nombre de liens de G_a et de $G_{a'}$.

Si $k = 2$, alors a et a' sont conclusions de deux liens affaiblissement et $G_a = G_{a'}$.

Autrement, si $A = B \otimes C$ (resp. $A = B \wp C$), alors soient b et b' les prémisses de type B et c et c' les prémisses de type C du lien $\otimes$ (resp. $\wp$) dont a et a' sont conclusions. Par définition $|e_n(a)| = (|e_n(b)|, |e_n(c)|)$ et $|e_n(a')| = (|e_n(b')|, |e_n(c')|)$. On aura $G_b \neq G_{b'}$ ou/et $G_c \neq G_{c'}$. Supposons par exemple que $G_b \neq G_{b'}$: par hypothèse d'induction on aura alors $|e_n(b)| \sim |e_n(b')|$. Dans le cas du $\otimes$, ceci suffit pour conclure que $|e_n(a)| \sim |e_n(a')|$. Dans le cas du $\wp$, remarquons que si $G_c = G_{c'}$ alors par ce qui précède $|e_n(c)| = |e_n(c')|$, tandis que si $G_c \neq G_{c'}$ alors par hypothèse d'induction $|e_n(c)| \sim |e_n(c')|$: quoiqu'il en soit on aura donc bien $|e_n(a)| \sim |e_n(a')|$.

Si $A = !B$, alors la conclusion est une application immédiate de l'hypothèse d'induction.

Si $A = ?B$, appelons m (resp. m') le lien dont a (resp. a') est conclusion. Supposons par exemple que $s(G_a) \geq s(G_{a'})$, et traitons d'abord le cas où $a' \in G_a$. On a $|e_n(a')| \subseteq |e_n(a)|$, et donc $|e_n(a')| \cup |e_n(a)|$ est une clique de $|?B| = B^\perp$: par conséquent $|e_n(a')| \preceq |e_n(a)| (?B)$. Puisque $a \neq a'$ et $n > 1$, $|e_n(a)| \neq |e_n(a')|$. Donc $|e_n(a')| \sim |e_n(a)|$.

On peut donc tourner notre attention vers le cas $a' \notin G_a$, c.à.d. $G_a \cap G_{a'} = \emptyset$. Si m ou m' (mais pas les deux !) est un lien affaiblissement, alors $|e_n(a')| \sim \emptyset$ ou bien $|e_n(a)| \sim \emptyset$: on peut donc supposer que m et m' sont des liens $?de$, pax ou $?co$.

Soient donc $a_1, \dots, a_l$ (resp. $a'_1, \dots, a'_q$) les prémisses des l (resp. q) liens dérélition au-dessus de a (resp. a'). Remarquons que puisque $G_a \cap G_{a'} = \emptyset$ on aura $a_i \neq a'_j \forall i \in \{1, \dots, l\}$ et $\forall j \in \{1, \dots, q\}$. Il n'y aura aucune arête de type atomique dans G_{a_i} ni dans $G_{a'_j} \forall i \in \{1, \dots, l\}$ et $\forall j \in \{1, \dots, q\}$ (puisque'il n'y en a pas dans G_a ni dans $G_{a'}$). $\forall i \in \{1, \dots, l\}$ et $\forall j \in \{1, \dots, q\}$ on a donc deux possibilités pour G_{a_i} et $G_{a'_j}$: soit $G_{a_i} = G_{a'_j}$ et alors $|e_n(a_i)| = |e_n(a'_j)|$, soit $G_{a_i} \neq G_{a'_j}$ et alors (par hypothèse d'induction) $|e_n(a_i)| \sim |e_n(a'_j)|$. Dans tous les cas on aura donc $|e_n(a')| \preceq |e_n(a)|$ (puisque $|e_n(a)| = \{|e_n(a_1)|, \dots, |e_n(a_l)|\}$ et $|e_n(a')| = \{|e_n(a'_1)|, \dots, |e_n(a'_q)|\}$).

Si $GC_a = GC_{a'}$ (cf. la définition 8.3.1 pour la notation), alors (puisque $G_a \neq G_{a'}$) il existe forcément $i \in \{1, \dots, l\}$ tel que $\forall j \in \{1, \dots, q\} G_{a_i} \neq G_{a'_j}$, ou bien il existe $j \in \{1, \dots, q\}$ tel que $\forall i \in \{1, \dots, l\} G_{a_i} \neq G_{a'_j}$: dans les deux cas l'hypothèse d'induction permet de conclure que $|e_n(a)| \neq |e_n(a')|$.

Si $GC_a \neq GC_{a'}$, alors par les lemmes 8.3.4 et 8.3.3 $\text{card}(|e_n(a)|) \neq \text{card}(|e_n(a')|)$ et donc $|e_n(a)| \neq |e_n(a')|$. (Remarquons qu'on utilise ici l'hypothèse $n > \sup(1, h)$, et c'est indispensable : sans cette hypothèse le lemme est faux). $\square$

11.1.4. LEMME. Soient a et a' deux arêtes différentes de même type A du réseau R , et e_1 une 1-expérience injective de R .

- (1) Si $a \in G_{a'}$ ou bien $a' \in G_a$, soit une parmi a et a' est conclusion d'un lien $?co$ et alors $e_1(a) \neq e_1(a')$, soit ce n'est pas le cas et alors $e(a) = e(a')$.
- (2) Si $G_a \cap G_{a'} = \emptyset$ et il existe une arête α de type atomique t.q. $\alpha \in G_a$ ou bien $\alpha \in G_{a'}$, alors $e_1(a) \neq e_1(a')$.
- (3) Si $G_a \cap G_{a'} = \emptyset$ et il n'y a pas d'arête de type atomique α t.q. $\alpha \in G_a$ ou bien $\alpha \in G_{a'}$, alors $e_1(a) \sim e_1(a')$.

Démonstration : (1) Le résultat est une conséquence du lemme 8.2.1 et de la définition de 1-expérience.

(2) Si $a \notin G_{a'}$ et $a' \notin G_a$, alors on procède exactement comme dans la démonstration de (1) du lemme 11.1.3.

(3) On peut procéder par induction sur $s(G_a) + s(G_{a'})$, comme dans la démonstration du lemme 11.1.3. $\square$

11.1.5. LEMME. Soient a et a' deux arêtes différentes de même type A du réseau R , et soit $n > \sup(1, h)$ où h est l'arité maximale des liens contraction de R . Soit e_1 une 1-expérience injective de R et e_n la n -expérience obsessionnelle de R induite par e_1 .

Si $e_1(a) \sim e_1(a')$, alors $|e_n(a)| \sim |e_n(a')|$.

Démonstration : On procède par induction sur $k = s(G_a) + s(G_{a'})$, la somme du nombre de liens de G_a et de $G_{a'}$.

Si $k = 0$, alors a et a' sont deux arêtes de type atomique et $|e_n(a)| = e_1(a)$, $|e_n(a')| = e_1(a')$.

Autrement, si $A = B \otimes C$ (resp. $A = B \wp C$), alors soient b et b' les prémisses de type B et c et c' les prémisses de type C des liens $\otimes$ (resp. $\wp$) dont a et a' sont conclusions. Par définition $|e_n(a)| = (|e_n(b)|, |e_n(c)|)$, $|e_n(a')| = (|e_n(b')|, |e_n(c')|)$, $e_1(a) = (e_1(b), e_1(c))$, $e_1(a') = (e_1(b'), e_1(c'))$. On aura $e_1(b) \sim e_1(b')$ ou/et $e_1(c) \sim e_1(c')$. Supposons par exemple que $e_1(b) \sim e_1(b')$: par hypothèse d'induction on aura alors $|e_n(b)| \sim |e_n(b')|$. Dans le cas du $\otimes$, ceci suffit pour conclure $|e_n(a)| \sim |e_n(a')|$. Dans le cas du $\wp$, si $e_1(c) \sim e_1(c')$, alors l'hypothèse d'induction permet de conclure. Mais il se pourrait que $e_1(c) = e_1(c')$. Dans ce cas, si il y avait une arête atomique dans G_c ou dans $G_{c'}$, alors (par le lemme 11.1.4) on aurait $c' \in G_c$ ou bien $c \in G_{c'}$, ce qui est impossible : il n'y a donc aucune arête de type atomique dans G_c et dans $G_{c'}$. Le lemme 11.1.3 nous dit que dans ce cas $|e_n(c)| \sim |e_n(c')|$, ce qui suffit pour conclure que $|e_n(a)| \sim |e_n(a')|$.

Si $A = !B$, alors la conclusion est une application immédiate de l'hypothèse d'induction.

Si $A = ?B$, appelons m (resp. m') le lien dont a (resp. a') est conclusion. Si m et m' sont deux liens affaiblissement, alors $e_1(a) = e_1(a')$. Si un exactement parmi m et m' est un lien affaiblissement, alors $|e_n(a)| \smile \emptyset = |e_n(a')|$ ou bien $|e_n(a')| \smile \emptyset = |e_n(a)|$: on peut donc supposer que m et m' sont des liens $?de$, pax ou $?co$.

Si $a \in G_{a'}$ ou bien $a' \in G_a$, alors on procède comme dans la démonstration du lemme 11.1.3 : $|e_n(a')| \smile |e_n(a)|$.

On va donc se placer sous l'hypothèse $G_a \cap G_{a'} = \emptyset$. Soient $a_1, \dots, a_l$ (resp. $a'_1, \dots, a'_q$) les prémisses des l (resp. q) liens dérélction au-dessus de a (resp. a'). Remarquons que puisque $G_a \cap G_{a'} = \emptyset$ on aura $G_{a_i} \cap G_{a'_j} = \emptyset \forall i \in \{1, \dots, l\}$ et $\forall j \in \{1, \dots, q\}$. De $e_1(a) \smile e_1(a')$ on déduit que $e_1(a_i) \smile e_1(a'_j) \forall i \in \{1, \dots, l\}$ et $\forall j \in \{1, \dots, q\}$. Si $e_1(a_i) \smile e_1(a'_j)$, alors par hypothèse d'induction $|e_n(a_i)| \smile |e_n(a'_j)|$. Si par contre $e_1(a_i) = e_1(a'_j)$, alors par le lemme 11.1.4 il n'y a pas d'arêtes de type atomique dans G_{a_i} ni dans $G_{a'_j}$ (puisque $G_a \cap G_{a'} = \emptyset$). Par le lemme 11.1.3 on aura alors $|e_n(a_i)| \smile |e_n(a'_j)|$. Tout ceci permet de conclure que $|e_n(a')| \smile |e_n(a)|$.

Si dans G_a et dans $G_{a'}$ il n'y a aucune arête atomique, alors (par le lemme 11.1.3) soit $|e_n(a')| \smile |e_n(a)|$ et on a fini, soit $|e_n(a')| = |e_n(a)|$ et alors $G_a = G_{a'}$. Mais dans ce dernier cas $e_1(a) = e_1(a')$, ce qui n'est pas vrai.

Traitons maintenant le cas où il existe une arête de type atomique dans G_a ou dans $G_{a'}$: appelons-la α et supposons (par exemple) que ce soit une arête de G_a . Soit $i \in \{1, \dots, l\}$ t.q. $\alpha \in G_{a_i}$. Par le lemme 11.1.4 on aura $\forall j \in \{1, \dots, q\}$ $e_1(a_i) \neq e_1(a'_j)$, donc $e_1(a_i) \smile e_1(a'_j)$. Par hypothèse d'induction on aura donc $\forall j \in \{1, \dots, q\}$ $|e_n(a_i)| \smile |e_n(a'_j)|$, et $|e_n(a)| \neq |e_n(a')|$ (dont on déduira $|e_n(a)| \smile |e_n(a')|$). $\square$

11.1.6. PROPOSITION. Soit R un réseau et $n > \sup(1, h)$ où h est l'arité maximale des liens contraction de R .

Si il existe une 1-expérience injective de R , alors il existe aussi une n -expérience obsessionnelle injective de R : il s'agit de (l'unique) n -expérience obsessionnelle de R induite par e_1 .

Démonstration : Par induction sur une séquentialisation π de R . Le cas le plus significatif est celui où la dernière règle de π est une règle r_m de contraction d'arité k . Soit alors π_1 la sous-preuve de π obtenue en effaçant r_m ainsi que sa conclusion et soit R_1 le sous-réseau de R obtenu en effaçant le dernier lien contraction m d'arité k correspondant à r_m . π_1 est une séquentialisation de R_1 , et la restriction e_1^1 de la 1-expérience injective e_1 de R à R_1 est une 1-expérience injective de R_1 . Soient $a_1, \dots, a_k$ les prémisses de m et a sa

conclusion. On a $e_1^1(a_i) \preceq e_1^1(a_j) \forall i, j \in \{1, \dots, k\}$. Soit e_n^1 la n -expérience obsessionnelle de R_1 fournie par l'hypothèse d'induction. Remarquons que puisque $\forall i \in \{1, \dots, k\}$ a_i est à profondeur 0 dans R_1 , $e_n^1(a_i)$ est un singleton.

Si pour $i, j \in \{1, \dots, k\}$ $e_1^1(a_i) = e_1^1(a_j)$, alors (a_i et a_j étant prémisses du même lien ?*co* de R on a $G_{a_i} \cap G_{a_j} = \emptyset$) par le lemme 11.1.4 il n'y a pas d'arêtes de type atomique dans G_{a_i} ni dans G_{a_j} : dans ce cas le lemme 11.1.3 nous permet d'affirmer que $|e_n^1(a_i)| \preceq |e_n^1(a_j)|$.

$\forall i, j \in \{1, \dots, k\}$ on a deux possibilités : soit $e_1^1(a_i) \sim e_1^1(a_j)$ et alors par le lemme 11.1.5 $|e_n^1(a_i)| \preceq |e_n^1(a_j)|$, soit $e_1^1(a_i) = e_1^1(a_j)$ et alors on vient de voir que $|e_n^1(a_i)| \preceq |e_n^1(a_j)|$. Dans tous les cas on a $|e_n^1(a_i)| \preceq |e_n^1(a_j)| \forall i, j \in \{1, \dots, k\}$. On peut donc définir l'expérience e_n de R t.q. $\forall a'$ arête de R , si $a' \neq a$ alors $e_n(a') = e_n^1(a')$, et $e_n(a) = \{|e_n^1(a_1)| \cup \dots \cup |e_n^1(a_k)|\}$.

Mentionnons aussi le cas où la dernière règle de π est une promotion. On procède de la même façon : l'expérience e_n cherchée est obtenue à partir de l'expérience e_n^1 fournie par l'hypothèse d'induction en "répétant n fois" e_n^1 , suivant la définition de n -expérience obsessionnelle (cf. remarque 8.1.2). $\square$

11.2 Réduction au cas sans boîtes

On a donc ramené le problème de l'existence d'une n -expérience obsessionnelle injective d'un réseau R à celle de l'existence d'une 1-expérience injective de R . On va maintenant ramener ce dernier problème à celui de l'existence d'une expérience injective pour le "linéarisé de R ", qui ne contiendra plus de boîtes.

11.2.1. REMARQUE. (et Définition). (i) A tout réseau R on peut associer de façon canonique un (unique) réseau sans boîtes que nous appellerons **le linéarisé** de R et que nous noterons $L(R)$: il s'agit du réseau obtenu en effaçant toutes les connexions entre les portes des boîtes, et tous les liens bien sûr et pax de R . Le fait que $L(R)$ soit un réseau standard est évident. On notera L l'application qui associe à toute formule A la formule $L(A)$, obtenue en effaçant toute occurrence du connecteur "!" dans A . Remarquons que les types des conclusions de R ne sont pas en règle générale les types des conclusions de $L(R)$: si les conclusions de R sont de type Γ , alors celles de $L(R)$ seront de type $L(\Gamma)$ (avec la convention usuelle pour ce genre de notation).

(ii) Soit R un réseau et $L(R)$ son linéarisé. Il existe une application canonique L qui associe à toute arête a de type A de R l'arête $L(a)$ de type $L(A)$ de

$L(R)$. Remarquons que L est loin d'être injective, si dans R il y a au moins une boîte.

(iii) Soit R un réseau et $L(R)$ son linéarisé. Si e_L est une expérience injective de $L(R)$, alors il existe au plus une 1-expérience e de R , telle que pour toute arête atomique α de R (donc de $L(R)$) $e(\alpha) = e_L(\alpha)$. On dira que e est la **délinéarisée** de e_L . Il est bien évident que si e existe elle est injective.

11.2.2. LEMME. Soit R un réseau, $L(R)$ son linéarisé, a et a' deux arêtes différentes de même type A de R . Soit e_L une expérience injective de $L(R)$ et e la délinéarisée de e_L .

Si $e_L(L(a)) \sim_{e_L(L(a'))}(\mathcal{L}(\mathcal{A}))$, alors $e(a) \sim e(a')(\mathcal{A})$.

Démonstration : On procède comme d'habitude par induction sur $s(G_a) + s(G_{a'})$.

Si $A = B \otimes C$ (resp. $A = B \wp C$), alors $L(A) = L(B) \otimes L(C)$ (resp. $L(A) = L(B) \wp L(C)$). Soient b et b' les prémisses de type B et c et c' les prémisses de type C du lien $\otimes$ (resp. $\wp$) dont a et a' sont conclusions. On a $e_L(L(a)) = (e_L(L(b)), e_L(L(c)))$. De $e_L(L(a)) \sim_{e_L(L(a'))}(\mathcal{L}(\mathcal{A}))$, on déduit $e_L(L(b)) \sim_{e_L(L(b'))}(\mathcal{L}(\mathcal{B}))$ ou bien $e_L(L(c)) \sim_{e_L(L(c'))}(\mathcal{L}(\mathcal{C}))$. Supposons par exemple que $e_L(L(b)) \sim_{e_L(L(b'))}(\mathcal{L}(\mathcal{B}))$: par hypothèse d'induction on aura alors $e(b) \sim e(b')(\mathcal{B})$. Dans le cas du $\otimes$, ceci suffit pour conclure $e(a) \sim e(a')(\mathcal{A})$. Dans le cas du $\wp$, si $e_L(L(c)) = e_L(L(c'))$, alors par le lemme 11.1.4 appliqué à l'expérience e_L de $L(R)$, $G_{L(c)}^{L(R)}$ et $G_{L(c')}^{L(R)}$ ne contiennent aucune arête atomique. Donc G_c^R et $G_{c'}^R$ ne contiennent aucune arête atomique, et par le lemme 11.1.4 on a $e(c) \sim e(c')(\mathcal{C})$, ce qui permet de conclure que $e(a) \sim e(a')(\mathcal{A})$.

Si $A = !B$, alors la conclusion est une application de l'hypothèse d'induction et de la définition de l'application L .

Si $A = ?B$, alors $L(A) = ?L(B)$. Si $a \in G_{a'}$ ou bien $a' \in G_a$, alors ni a ni a' ne peut être conclusion d'un lien affaiblissement. Supposons par exemple que $a' \in G_a$. On aura $e(a') \subseteq e(a)$, donc $e(a) \sim e(a')(\mathcal{A})$: reste à démontrer que $e(a) \neq e(a')$. Puisque $a' \in G_a$, $L(a') \in G_{L(a)}$ (avec, éventuellement, $L(a) = L(a')$). $e_L(L(a)) \sim_{e_L(L(a'))}(\mathcal{L}(\mathcal{A}))$ implique $L(a) \neq L(a')$: puisque dans $L(R)$ il n'y a pas de liens pax, la seule possibilité est que $L(a')$ soit conclusion d'un lien $?de$ au-dessus de $L(a)$ qui elle est conclusion d'un lien $?co$. Par le lemme 8.2.1 on aura alors $e(a) \neq e(a')$.

Supposons maintenant que $G_a \cap G_{a'} = \emptyset$. Puisque $e_L(L(a)) \neq e_L(L(a'))$ une parmi les arêtes a et a' n'est pas conclusion d'un lien affaiblissement. Si (par exemple) a est conclusion d'un lien affaiblissement, alors a' ne l'est pas et

par le lemme 8.2.1 $e(a) \neq e(a')$: donc $\emptyset = e(a) \sim e(a')$. Appelons m (resp. m') le lien dont a (resp. a') est conclusion : on peut supposer que m et m' sont des liens $?de$, pax ou $?co$.

Soient $a_1, \dots, a_l$ (resp. $a'_1, \dots, a'_q$) les prémisses des l (resp. q) liens dérélction au-dessus de a (resp. a'). Remarquons que puisque $G_a \cap G_{a'} = \emptyset$ on aura $G_{a_i} \cap G_{a'_j} = \emptyset \forall i \in \{1, \dots, l\}$ et $\forall j \in \{1, \dots, q\}$. On a :

$e_L(L(a)) = \{e_L(L(a_1)), \dots, e_L(L(a_l))\}$ et $e(a) = \{e(a_1), \dots, e(a_l)\}$,

$e_L(L(a')) = \{e_L(L(a'_1)), \dots, e_L(L(a'_q))\}$ et $e(a') = \{e(a'_1), \dots, e(a'_q)\}$.

De $e_L(L(a)) \sim_{e_L(L(a'))}(\mathcal{L}(\mathcal{A}))$, on déduit que $\forall i \in \{1, \dots, l\}$ et $\forall j \in \{1, \dots, q\}$

$e_L(L(a_i)) \sim e_L(L(a'_j))$. Si $e_L(L(a_i)) \sim e_L(L(a'_j))$, alors par hypothèse d'induction $e(a_i) \sim e(a'_j)$. Si par contre $e_L(L(a_i)) = e_L(L(a'_j))$, alors par le lemme 11.1.4 il n'y a pas d'arêtes de type atomique dans $G_{L(a_i)}$ ni dans $G_{L(a'_j)}$

(puisque $G_a \cap G_{a'} = \emptyset$, et donc $G_{L(a)} \cap G_{L(a')} = \emptyset$). Donc il n'y a pas d'arêtes de type atomique dans G_{a_i} ni dans $G_{a'_j}$: toujours par le lemme 11.1.4 on

aura alors $e(a_i) \sim e(a'_j)$. Tout ceci permet de conclure que $e(a) \sim e(a')(\mathcal{A})$.

De $e_L(L(a)) \sim_{e_L(L(a'))}(\mathcal{L}(\mathcal{A}))$, on déduit aussi que $card(e_L(L(a))) \neq card(e_L(L(a')))$

(c.à.d. $l \neq q$), ou bien l'existence de $i \in \{1, \dots, l\}$ t.q. $\forall j \in \{1, \dots, q\}$

$e_L(L(a_i)) \sim e_L(L(a'_j))$, ou bien l'existence de $j \in \{1, \dots, q\}$ t.q. $\forall i \in \{1, \dots, l\}$

$e_L(L(a_i)) \sim e_L(L(a'_j))$. Dans les deux derniers cas on appliquera l'hypothèse d'induction et on en déduira que $e(a) \sim e(a')$. Dans le premier cas, on aura

aussi $e(a) \sim e(a')$. $\square$

11.2.3. PROPOSITION. Soit R un réseau et $L(R)$ son linéarisé. Si il existe une expérience injective e_L de $L(R)$, alors il existe une 1-expérience injective e de R : il s'agit de la délinéarisée de e_L .

Démonstration : On procède comme dans la démonstration de la proposition 11.1.6 par induction sur une séquentialisation π de R . Le cas plus significatif sera encore une fois celui où la dernière règle de π est une règle de contraction, et on appliquera cette fois le lemme 11.2.2. $\square$

11.2.4. PROPOSITION. Soit R un réseau, $L(R)$ son linéarisé et $n > \sup(1, h)$ où h est l'arité maximale des liens contraction de R .

Si il existe une expérience injective e_L de $L(R)$, alors il existe une n -expérience obsessionnelle injective de R .

Démonstration : Conséquence immédiate des propositions 11.1.6 et 11.2.3.

$\square$

Chapitre 12

Elimination des liens par

Nous avons ramené (grâce à la proposition 11.2.4) la question de l'existence d'une expérience obsessionnelle et injective d'un réseau R , à celle de l'existence d'une expérience injective du réseau sans boîtes $L(R)$. Dans le même ordre d'idées, nous allons à présent ramener (en l'absence d'affaiblissements) cette dernière question à celle de l'existence d'une expérience injective d'un réseau quelconque de $L(R)^{\wp}$ (défini dans la remarque 12.0.9), qui ne contient (aucune boîte et) aucun lien $\wp$.

Nous définissons une procédure pour “éliminer les liens $\wp$ d'un réseau sans boîtes”, et nous montrons que si il existe une expérience injective d'un réseau R^- obtenu à partir de R (sans boîtes) en appliquant cette procédure, alors il existe aussi une expérience injective de R (proposition 12.0.10).

Dans ce chapitre, tout réseau sera un réseau (standard) ne contenant aucune boîte, et ne contenant aucun lien affaiblissement.

Soit R un réseau. Rappelons que puisque R est sans boîtes, une expérience de R associe à toute arête de R une unique étiquette : c'est un étiquetage des arêtes de R (comme dans le cas des réseaux multiplicatifs).

Si à tout lien axiome l de conclusion α_l et $\alpha_l^\perp$ on associe un élément x_l de la trame de l'espace cohérent $\mathcal{A}$ de telle sorte que si $l \neq l'$ alors $x_l \neq x_{l'}$, qu'est-ce qui pourrait s'opposer, à présent, à l'existence d'une expérience e de R t.q. pour tout lien axiome l on ait $e(\alpha_l) = e(\alpha_l^\perp) = x_l$? Seulement la présence des liens ?*co*.

12.0.5. REMARQUE. Pour que l'étiquetage des arêtes de R induit par la définition ci-dessus donne lieu à une expérience de R *il faut et il suffit* que

pour tout lien $?co$ de R de prémisses $b_1, \dots, b_k$ on ait $e(b_i) \prec e(b_j)$, et ce $\forall i, j \in \{1, \dots, k\}$.

12.0.6 Conventions :

Nous dirons dans la suite qu'une arête a est prémisses d'un lien n (resp. terminale) "modulo un lien m ", lorsque a est prémisses de n (resp. terminale) ou bien a est une prémisses de m et la conclusion de m est prémisses de n (resp. terminale).

Nous dirons d'un lien $?co$ (ou $?w$ dans les chapitres suivants) qu'il est terminal dans un réseau R lorsque la conclusion du lien en question est aussi une conclusion du réseau R .

12.0.7 La procédure d'élimination des liens par

Soit R un réseau et Γ le séquent conclusion de R . Soit $A \in \Gamma$ t.q. $C \wp D$ est une *occurrence* de sous-formule de A .

Soient $a_1, \dots, a_k$ les arêtes de G_a de type $C \wp D$ (nous parlons toujours de l'*occurrence* de sous-formule $C \wp D$ de A). Soient $n_1, \dots, n_k$ les k liens $\wp$ de conclusions respectives $a_1, \dots, a_k$, soient $c_1, \dots, c_k$ (resp. $d_1, \dots, d_k$) les prémisses de type C (resp. D) de $n_1, \dots, n_k$.

Appelons R' le graphe obtenu à partir de R de la façon suivante.

Pour tout $i \in \{1, \dots, k\}$, on effectue les trois premières opérations, et on complète la transformation du graphe R en R' par la quatrième :

1. on efface le lien n_i et sa conclusion a_i
2. si a_i est prémisses d'un lien m_i , alors on substitue la prémisses a_i de m_i par d_i
3. si la formule C ne commence pas par un "?", alors on rajoute un lien déréluction de prémisses c_i et de conclusion une arête g'_i de type $?C$ (autrement $g'_i = c_i$ sera de type C)
4. Si $k \geq 2$, alors on rajoute au graphe obtenu après avoir appliqué les opérations 1-3, un lien contraction de prémisses $g'_1, \dots, g'_k$ et de conclusion l'arête g (de type $?C$ ou bien C).

Le lecteur n'aura aucun mal à se convaincre que le graphe ainsi obtenu est un réseau standard dont le séquent conclusion est $\Gamma \setminus A, A[D/C \wp D], ?C$ (ou bien $\Gamma \setminus A, A[D/C \wp D], C$). Il faut simplement remarquer que, si jamais certaines arêtes parmi les c_i sont conclusions d'un lien $?co$, alors le lien $?co$ de R' de conclusion g aura plus de k prémisses.

En inversant le rôle des arêtes c_i et d_i , il est bien évident qu'on peut obtenir un réseau dont le séquent conclusion est $\Gamma \setminus A, A[C/C \wp D], ?D$ (ou bien

$\Gamma \setminus A, A[C/C \wp D], D).$

On dira que le réseau R' est obtenu à partir de R en éliminant un $\wp$.

12.0.8. LEMME. Soit R un réseau et R' un réseau obtenu à partir de R en éliminant un $\wp$.

Si il existe une expérience injective e' de R' , alors il existe une expérience injective e de R .

Démonstration : Nous reprenons, dans cette démonstration, les mêmes notations que dans la définition de la procédure d'élimination des $\wp$. En particulier, on supposera que le séquent conclusion de R (resp. R') est Γ (resp. $\Gamma \setminus A, A[D/C \wp D], ?C$, ou bien $\Gamma \setminus A, A[D/C \wp D], C$).

Toute arête b de R différente de $a_1, \dots, a_k$ est une arête de R' que nous noterons b' . *Attention !* Les arêtes b et b' ne sont pas forcément de même type : si on note B le type de l'arête b' de R' , alors l'arête b de R est soit de type B soit de type $B[C \wp D/D]$.

Nous allons montrer que l'expérience e de R cherchée n'est autre que "l'expérience e' définie sur les arêtes de R de la seule façon possible".

Pour définir correctement l'expérience e , nous introduisons une terminologie dont on se servira seulement dans cette démonstration. On dira qu'une arête b de R t.q. $b \notin \{a_1, \dots, a_k\}$ est "spéciale", lorsqu'il existe un chemin (il s'agit en fait d'un chemin direct, suivant la définition 2.1.2) contenant b , et dont la première arête est l'une des a_i ($i \in \{1, \dots, k\}$) et la dernière arête est une conclusion de R . Remarquons que toute arête spéciale b de R est une arête b' de R' , qui n'a pas le même type que b .

Pour toute arête b de R :

- si b n'est pas spéciale et $b \notin \{a_1, \dots, a_k\}$, alors on pose $e(b) = e'(b')$
- si $b = a_i$ pour un certain $i \in \{1, \dots, k\}$, alors on pose $e(a_i) = (e'(c'_i), e'(d'_i))$
- si b est une arête spéciale, alors la définition de $e(b)$ découle de celles que nous venons de donner (en suivant la définition 6.2.1).

Ce qu'il faut prouver c'est que l'étiquetage de R ainsi défini est bien une expérience de R . On a vu dans la remarque 12.0.5 que la seule chose qui pourrait s'y opposer c'est la présence d'un lien contraction m de R de conclusion l'arête q et ayant comme prémisses les arêtes $q_1, \dots, q_h$, telles que $e(q_l) \frown e(q_s)$ pour certains $l, s \in \{1, \dots, h\}$. Si aucune des arêtes $q_1, \dots, q_h$ n'est spéciale, alors puisque pour tout $l, s \in \{1, \dots, h\}$ on a $e'(q'_l) \frown e'(q'_s)$, on aura (par définition de e), $e(q_l) \frown e(q_s)$. Autrement, *toutes* les arêtes $q_1, \dots, q_h$ (ainsi que q) sont spéciales : il suffira donc de montrer que dans ce cas $\forall l, s \in \{1, \dots, h\}$ on a $e(q_l) \frown e(q_s)$. Remarquons que les arêtes $q'_1, \dots, q'_h$ sont prémisses d'un lien $?co$ de R' .

Nous allons prouver que $\forall i, j \in \{1, \dots, k\}$, on a la propriété $(\star)$ suivante :

- ($\star.1$) si $e'(d'_i) \smile e'(d'_j)$, alors $e(a_i) \smile e(a_j)$
- ($\star.2$) si $e'(d'_i) = e'(d'_j)$, alors $e(a_i) \smile e(a_j)$.

Montrons tout d'abord que ceci implique que $\forall l, s \in \{1, \dots, h\}$ on a $e(q_l) \smile e(q_s)$ (le lemme est alors démontré puisque l'unique cas problématique, évoqué ci-dessus, ne peut pas se produire).

Fixons deux arêtes distinctes t_1 et t_2 de même type T , spéciales, et telles que $G_{t_1} \cap G_{t_2} = \emptyset$. Intuitivement, $(\star)$ affirme que si $e'(t'_1) \smile e'(t'_2)$ et si jamais la relation de cohérence entre $e'(d'_i)$ et $e'(d'_j)$ joue un rôle quelconque dans le fait que $e'(t'_1) \smile e'(t'_2)$, ce même rôle sera joué par la relation de cohérence entre $e(a_i)$ et $e(a_j)$, et on pourra donc affirmer que $e(t_1) \smile e(t_2)$. Plus précisément, nous allons montrer que :

- si $e'(t'_1) \smile e'(t'_2)$, alors $e(t_1) \smile e(t_2)$
- si $e'(t'_1) = e'(t'_2)$, alors $e(t_1) \smile e(t_2)$.

Du fait que $e'(q'_l) \smile e'(q'_s)$, on pourra alors déduire que $e(q_l) \smile e(q_s)$.

Ce qu'on veut prouver est en fait assez clair, mais pour en donner une preuve (un peu plus) rigoureuse nous ne voyons d'autre moyen que celui de procéder par induction. Pour toute arête b de R , notons $\sharp\Phi_b$ le nombre d'arêtes du chemin dont b est la première arête et ayant une conclusion de R comme arête terminale. On peut faire une preuve par induction sur $p = \sum_{i=1}^k \sharp\Phi_{a_i} - (\sharp\Phi_{t_1} + \sharp\Phi_{t_2})$.

Reste donc à démontrer $(\star)$. Rappelons que $\forall i \in \{1, \dots, k\}$ les arêtes c_i et d_i (prémisses du lien n_i de conclusion a_i) ne sont pas spéciales, ce qui veut dire que $e(c_i) = e'(c'_i)$ et $e(d_i) = e'(d'_i)$.

La propriété $(\star)$ est en fait une conséquence du fait que $\forall i, j \in \{1, \dots, k\}$ les arêtes c'_i et c'_j de R' sont prémisses (modulo un éventuel lien $?de$) d'un lien $?co$: on a $e'(c'_i) \smile e'(c'_j)$ et donc $e(c_i) \smile e(c_j)$.

Soient alors $i, j \in \{1, \dots, k\}$. Par définition de la relation de cohérence dans l'espace cohérent $\mathcal{C} \wp \mathcal{D}$, on a :

- ($\star.1$) si $e'(d'_i) \smile e'(d'_j)$, alors $e(d_i) \smile e(d_j)$ et $e(c_i) \smile e(c_j)$: donc $e(a_i) \smile e(a_j)$.
- ($\star.2$) si $e'(d'_i) = e'(d'_j)$, alors $e(d_i) = e(d_j)$ et $e(c_i) \smile e(c_j)$: donc $e(a_i) \smile e(a_j)$.

□

12.0.9. REMARQUE. En appliquant le nombre de fois nécessaire la procédure d'élimination des $\wp$ à un réseau R , on obtiendra un réseau R^- qui ne contiendra plus aucun lien $\wp$: il n'y aura plus, dans R^- , que des liens ax (de conclusions atomiques), $?de$, $?co$, $\otimes$.

Bien sûr, R^- n'est pas unique : nous noterons dans la suite $R^{\wp}$ l'ensemble

des réseaux obtenus en éliminant tous les liens $\wp$, suivant la procédure décrite ci-dessus.

Il est important de remarquer que les seuls liens ?co des réseaux de $R^\wp$ qui ne sont pas des liens de R sont terminaux.

12.0.10. PROPOSITION. Soit R un réseau. Si il existe une expérience injective d'un réseau R^- de $R^\wp$, alors il existe une expérience injective de R .

Démonstration : Par induction sur le nombre de liens $\wp$ de R . Dans l'étape d'induction, on utilisera le lemme 12.0.8. $\square$

Chapitre 13

La sémantique cohérente n'est pas injective

Nous montrons qu'il n'existe pas (en général) d'expérience injective pour un réseau sans boîtes et sans liens $\emptyset$, et nous en déduisons *ipso facto* un contre-exemple à la conjecture 7.1.3 (pour la sémantique cohérente), ce qui montre bien la pertinence de notre approche.

Nous exhibons ensuite un autre contre-exemple à la conjecture 7.1.3 (toujours pour la sémantique cohérente), de nature un peu différente.

Quelques remarques sur la non injectivité de la sémantique cohérente s'imposent alors.

13.1 Le premier contre-exemple

La morale du chapitre 11 est que pour montrer que si R et R' sont deux réseaux sémantiquement équivalents alors $SPL(R) = SPL(R')$ ($R = R'$ en l'absence d'affaiblissements), il suffit de montrer l'existence d'une expérience injective pour tout réseau sans boîtes. Mais il s'avère que ceci est faux (de même que la conjecture 7.1.3), comme nous allons le voir dans le contre-exemple qui suit.

Il n'existe aucune expérience injective pour le réseau R de la figure 11.

Les chiffres que nous avons associés aux différentes arêtes indiquent quelles sont les contraintes que les liens contraction de R imposent : si aux deux arêtes a et a' de R est associé le même entier, alors toute expérience e de R doit satisfaire $e(a) \preceq e(a')$.

Soient $x, y \in |\mathcal{X}|$ les étiquettes que l'expérience e de R associe aux conclusions des deux liens axiome de R . On voit bien qu'il faut d'une part que

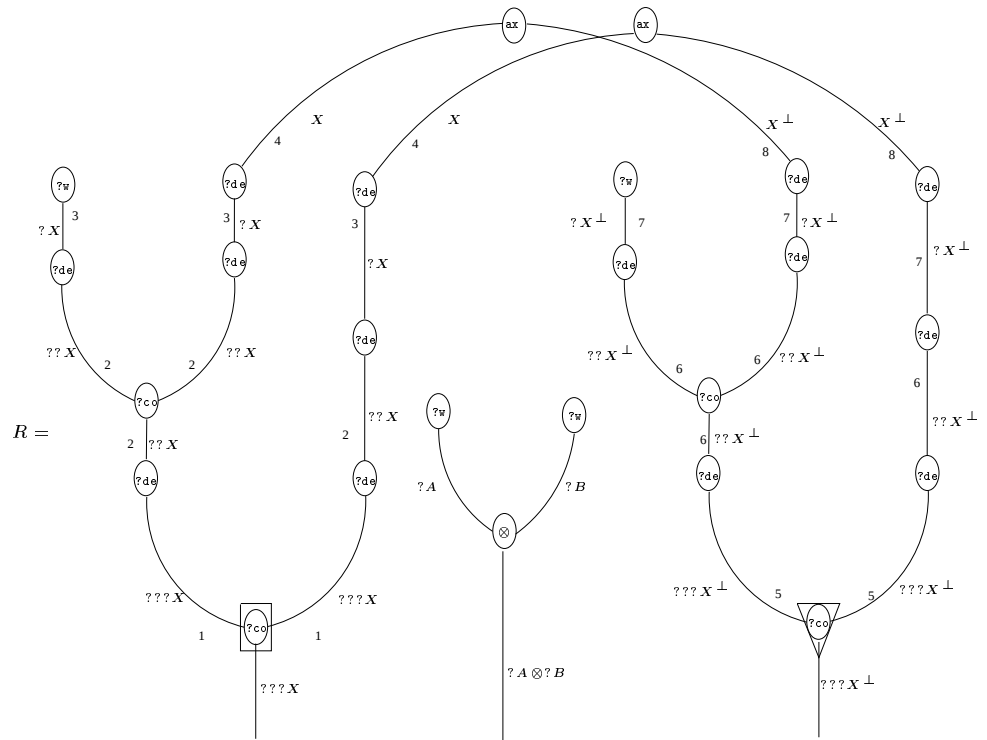


Figure 11. Il n'existe aucune expérience injective du réseau R .

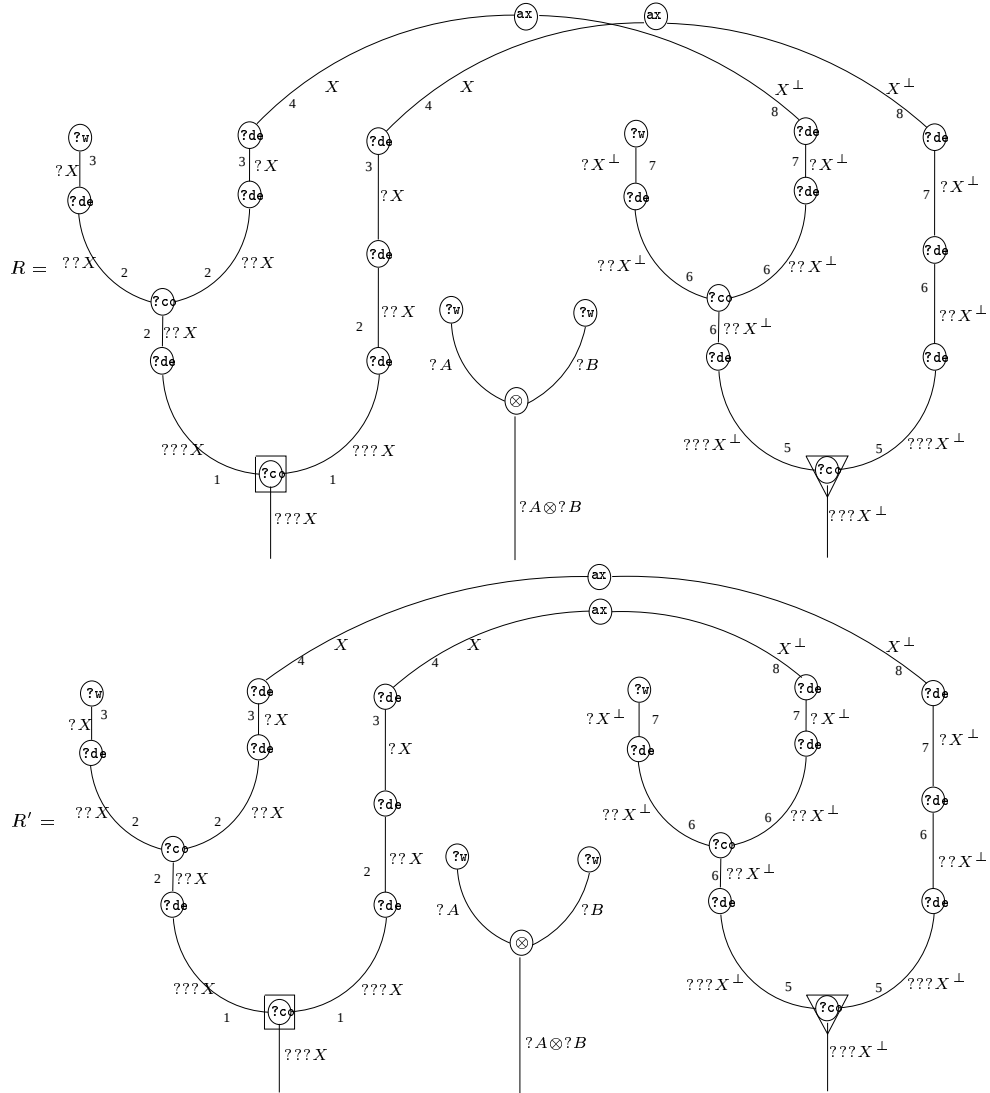


Figure 12. Contre-exemple 1 : pour les réseaux R et R' ci-dessus, on a $[R]_{coh} = [R']_{coh}$ et pourtant $R \neq R'$.

$x \widetilde{\sim} y(\mathcal{X})$ (à cause du lien $?co$ carré dans la figure) et d'autre part que $x \widetilde{\sim} y(\mathcal{X}^\perp)$ (à cause du lien $?co$ triangulaire dans la figure), c.à.d. que $x = y$. Il n'existe donc aucune expérience injective de R .

Le lecteur attentif, après avoir vu le réseau R , nous objectera qu'on peut faire plus simple, et il aura raison. Mais cet exemple n'a pas été choisi au hasard : si la sémantique ne peut pas distinguer les deux liens axiome de R , alors on peut “croiser les fils”, obtenir un réseau R' de même sémantique que R et qui est pourtant différent... Comme dans la figure 12 ! C'est pour que le réseau R soit différent de R' (et pour qu'ils soient tous les deux standard) que le réseau précédent a été choisi.

Pour se convaincre que $R \neq R'$ le lecteur pourra montrer (par exemple) qu'il existe un sous-réseau de R qui n'est pas un sous-réseau de R' . Le fait que R et R' soient sémantiquement équivalents (tant pour la sémantique ensembliste que pour la sémantique multiensembliste) est une conséquence immédiate de la non existence, pour R et R' , d'une expérience injective. Remarquons que bien-sûr $SPLO(R) = SPLO(R')$, en accord avec le théorème 10.2.3.

Nous ferons dans la suite de plus amples considérations au sujet de ce contre-exemple et du suivant. Bornons-nous ici à une simple remarque. Si on “ouvre”, dans R et R' , la contraction triangulaire ou la contraction carrée (c.à.d. qu'on efface le lien contraction et sa conclusion), on obtient deux réseaux R_1 et R'_1 de mêmes conclusions, qui sont sous-réseaux de R et R' respectivement. R_1 et R'_1 n'auront pas la même sémantique : la condition d'uniformité imposée par le lien contraction qu'on a effacé ayant disparu, il existe à présent des expériences injectives, tant de R_1 que de R'_1 , qui permettent de distinguer sémantiquement ces deux réseaux. Ceci semble, notamment, plaider en faveur d'une approche “non inductive” (telle que la notre) à la question de l'injectivité.

13.2 Le deuxième contre-exemple

Nous allons maintenant montrer qu'il existe un autre phénomène qui, lui aussi, conduit à un contre-exemple à la conjecture 7.1.3. En présence d'affaiblissements, la caractérisation des boîtes que fournit la proposition 10.0.7 est évidemment fausse. Nous allons montrer que (toujours à cause des requêtes d'uniformité) la sémantique d'un réseau est incapable, à elle seule, de reconstituer de façon unique les connexions entre les différentes portes des boîtes du réseau. Ce même contre-exemple montrera que même pour le système *ELL* défini dans [Girard 95b] la sémantique cohérente (ensembliste ou multiensembliste) n'est pas injective.

Les deux réseaux R et R' de la figure 13 sont (clairement) différents et ont pourtant la même sémantique cohérente (qu'il s'agisse de sémantique ensembliste ou de sémantique multiensembliste). Montrons que R et R' sont sémantiquement équivalents. Nous avons associé à certaines arêtes des chiffres, en suivant la même convention que pour le contre-exemple précédent : si aux deux arêtes a_1 et a_2 de R (resp. de R') est associé le même entier, alors toute expérience e de R (resp. de R') doit satisfaire $e(a_1) \preceq e(a_2)$. Remarquons que cette notation a un sens, puisque toutes les arêtes a auxquelles nous avons associé un entier sont à profondeur zéro, et $e(a)$ est donc un singleton (suivant la définition 6.2) : nous confondons ici, comme nous l'avons fait pour les 1-expériences, $e(a)$ avec l'unique élément que ce multiensemble contient.

Soient c_1 et c_2 (resp. c'_1 et c'_2) les conclusions des deux pal de profondeur zéro dans R (resp. dans R'). La condition d'uniformité provenant du lien $?co$ de profondeur zéro dans R (resp. dans R') impose que pour toute expérience e de R (resp. e' de R') $e(c_1) \preceq e(c_2)$ (resp. $e'(c'_1) \preceq e'(c'_2)$). Mais tout élément de $e(c_i)$ (resp. de $e'(c'_i)$), pour $i \in \{1, 2\}$, est de la forme $\{n_i[\emptyset]\}$ (resp. $\{n'_i[\emptyset]\}$) : la seule possibilité est alors que $n_1 = n_2$ (resp. $n'_1 = n'_2$). Dans le cas ensembliste on aura carément $n_1 = n_2 = n'_1 = n'_2 = 1$.

Il suffit alors de remarquer que le seul moyen pour la sémantique de distinguer R et R' est de réussir à exprimer le fait que le sous-graphe T de R et R' se trouve dans une boîte et pas dans l'autre : c'est précisément ce que l'uniformité interdit. La sémantique est donc incapable de nous dire dans quelle boîte se trouve le sous-graphe T de R et R' , ce qui permet de conclure : à toute expérience de R on peut associer une expérience de R' de même résultat (et réciproquement, bien sûr).

Nous n'avons pas du tout parlé, dans cette thèse, des éléments neutres de LL (voir aussi le paragraphe 15.2), mais une remarque s'impose ici : le deuxième contre-exemple est aussi un contre-exemple à l'injectivité de la sémantique cohérente *multiensembliste*, pour le fragment multiplicatif et exponentiel sans liens axiome mais avec les liens introduisant les constantes multiplicatives 1 et $\perp$. Dans le cas ensembliste, la non injectivité est triviale (cf. remarque 13.3.2).

On va conclure ce paragraphe en expliquant rapidement comment modifier les réseaux R et R' de telle sorte que le contre-exemple ci-dessus soit aussi un contre-exemple à l'injectivité pour *ELL*. Suivant la caractérisation de *ELL* de [DJ 99], un réseau R de *MELL* ne contenant que des liens axiome de conclusions atomiques (mais pas forcément standard) est un réseau de *ELL*

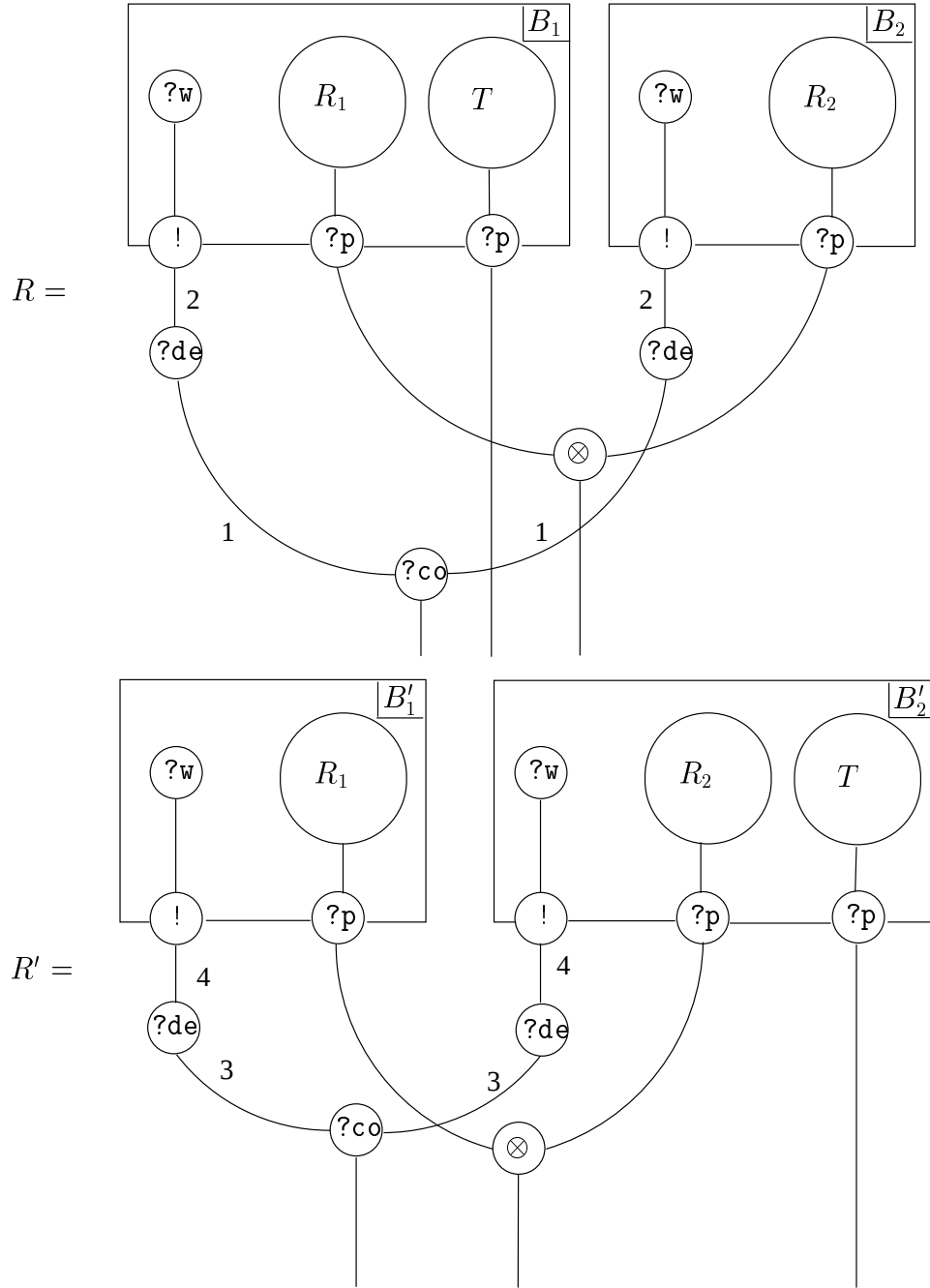


Figure 13. Contre-exemple 2 : pour les réseaux R et R' ci-dessus, on a $[R]_{coh} = [R']_{coh}$ et pourtant $R \neq R'$.

ssi le chemin direct issu de tout lien $?de$ de R qui a comme arête terminale une conclusion de R ou bien une prémisses d'un lien logique ou coupure, contient *exactement* un lien pax.

Les réseaux R et R' du deuxième contre-exemple ne sont pas des réseaux de ELL , mais on peut rajouter “les boîtes qu'il faut” pour qu'ils le deviennent (suivant la caractérisation ci-dessus) : il faut rajouter, dans R comme dans R' , une boîte dont la prémisses de la pal est la conclusion du lien pax dont la prémisses est l'arête conclusion de T . Il faudra aussi rajouter un lien $?de$ (de prémisses la conclusion du lien $\otimes$), et bien sûr choisir R_1 , R_2 et T de façon à ce que la condition ci-dessus soit satisfaite.

Ces deux réseaux construits (remarquons au passage qu'ils ne sont pas standard), nous n'avons plus qu'à vérifier qu'ils ont même sémantique : ceci est évident, puisque c'est le cas pour R et R' .

L'intérêt de la non injectivité de la sémantique cohérente pour ELL est limité, en ce sens que toute sémantique de LL est bien une sémantique de ELL , mais à ce jour il n'existe pas, à notre connaissance, de sémantique dénotationnelle propre à ELL .

13.3 Remarques au sujet de la non injectivité

La proposition suivante met en relation les deux contre-exemples ci-dessus avec des approches plus traditionnelles, comme celles de [Statman 83] ou de [Joly 00] au λ -calcul simplement typé. Ce n'est pas tout à fait notre point de vue, mais nous signalons quand même le résultat qui suit comme une conséquence de notre analyse.

13.3.1. PROPOSITION. Sur l'ensemble des réseaux de $MELL$ (pas forcément standard) il existe une relation d'équivalence $\sim$ t.q. si $R \sim R'$ alors R et R' sont deux réseaux de même conclusion, qui satisfait les conditions suivantes :

(1) elle est non dégénérée : il existe deux réseaux de même conclusion qui ne sont pas équivalents

(2) elle étend $\simeq_{\beta\eta}$: si $R \simeq_{\beta\eta} R'$, alors $R \sim R'$

(3) c'est une congruence : si $R \sim R'$, si R_1 (resp. R'_1) est un sous-réseau de R (resp. de R') dont les conclusions sont typées par Γ t.q. $R_1 \sim R'_1$, et si R_2 (resp. R'_2) est un réseau dont les conclusions sont typées par Γ t.q. $R_2 \sim R'_2$, alors $R[R_2/R_1] \sim R'[R'_2/R'_1]$

(4) elle satisfait “l'ambiguïté des types” : si $R \sim R'$ et si on substitue, dans R comme dans R' , toutes les occurrences de la variable propositionnelle X par une formule A de $MELL$ qui ne contient aucune occurrence de variable

propositionnelle apparaissant ailleurs dans R ou dans R' , alors $R[A/X] \sim R'[A/X]$, où $R[A/X]$ (resp. $R'[A/X]$) est obtenu à partir de R (resp. R') en effectuant les substitutions indiquées ci-dessus

(5) elle étend *strictement* $\simeq_{\beta\eta}$: il existe deux réseaux R et R' de $MELL$ t.q. $R \sim R'$ sans que $R \simeq_{\beta\eta} R'$.

Démonstration : Nous allons prouver que la relation d'équivalence cherchée est la relation d'équivalence sémantique (cf. 7.1.2), et ce tant pour la sémantique cohérente ensembliste que pour la sémantique cohérente multiensembliste. Notons $\sim_c$ une quelconque de ces deux relations d'équivalence.

$\sim_c$ satisfait clairement (2) et (3). Pour nous convaincre qu'elle satisfait (1), considérons un séquent prouvable Γ de MLL et l'union disjointe des arbres des formules de Γ : appelons-la S . Tout réseau standard R dont les conclusions sont de type Γ va satisfaire $SPO(R) = S$. Il est bien clair qu'on peut choisir Γ de telle sorte qu'en "croisant les fils des arêtes atomiques de S " autrement que comme indiqué par R on obtienne un réseau $R' \neq R$. On aura alors (par le théorème 7.3.4) $[R]_c \neq [R']_c$.

Les deux contre-exemples qui précèdent montrent bien que $\sim_c$ satisfait (5). Quant à (4), c'est une conséquence directe de la définition 7.1.2 : il y a une quantification universelle sur l'interprétation des formules atomiques. $\square$

13.3.2. REMARQUE. La proposition précédente reste bien sûr vraie si on élimine la condition (4), et la preuve est très simple : il suffit d'interpréter tous les atomes de R et R' par le même espace cohérent *fini* et considérer la sémantique *ensembliste*. On décrètera ensuite que $R \sim R'$ ssi l'interprétation de R et celle de R' sont égales. Puisqu'on peut construire une infinité de réseaux qui ne sont pas $\beta\eta$ -équivalents (deux à deux) tout en ayant le même séquent conclusion, et puisque l'interprétation de ce séquent est un espace cohérent *fini*, il est bien clair que plusieurs parmi ces réseaux seront identifiés par notre interprétation. Le lecteur qui désire se convaincre de l'existence des réseaux dont il est question ci-dessus pourra (par exemple) traduire en logique linéaire la formule intuitionniste $\forall X(X, (X \rightarrow X) \rightarrow X)$ qui est le type des entiers de Church, et prouver dans LL que $1, 2, \dots$ sont des entiers (en se bornant à écrire la partie propositionnelle de ces preuves). Cette même remarque s'applique aussi au fragment multiplicatif et exponentiel de LL, sans atomes mais avec les éléments neutres 1 et $\perp$.

13.3.3. REMARQUE. Une conséquence de la proposition précédente est que le résultat de Statman sur la maximalité de la $\beta\eta$ -équivalence (modulo l'ambiguïté des types) en λ -calcul simplement typé ne s'étend pas à $MELL$.

Il était tentant de conclure la remarque précédente en rappelant que les termes du λ -calcul simplement typé sont les preuves de la logique minimale, tandis que *MELL* permet de coder toute la logique classique. Mais fallait-il vraiment céder à cette tentation ? A notre avis, la réponse est négative (c'est pourquoi nous ne l'avons pas fait), et nous essayerons d'en convaincre le lecteur dans les chapitres qui suivent.

Nous aurons l'occasion de rediscuter des deux contre-exemples du paragraphe précédent. Mais on peut se poser dès maintenant une question naïve très naturelle : si syntaxe et sémantique ne sont pas d'accord, qui a raison ? En d'autres termes, est-ce “juste” d'identifier les deux réseaux R et R' des contre-exemples 1 et 2 ?

Il est bien évident que la réponse dépendra éminemment du système judiciaire dans lequel on se place... Mais si on prend comme critère (un peu informel) celui selon lequel il faudrait identifier toutes les preuves qui interagissent de la même façon avec l'environnement, i.e. avec les autres preuves en effectuant des coupures (en première approximation), on aurait plutôt tendance à répondre de façon affirmative à la question ci-dessus. En effet, le lecteur n'aura pas de mal à vérifier que les réseaux R et R' (tant dans le cas du contre-exemple 1 que dans celui du contre-exemple 2) auront le même comportement vis à vis de tout ensemble de réseaux avec lesquels on puisse couper R et R' . Ceci est essentiellement dû au fait que les prémisses d'un lien contraction ne pourront être coupées *que* sur une *même* boîte : c'est précisément ce “même” dont la sémantique cohérente est au courant, qu'elle traduit par les conditions que nous avons appelées “d'uniformité” (ce que nous venons de dire expliquant aussi, en partie, la terminologie). De ce point de vue, ce serait la sémantique qui “aurait raison”.

Ce n'est pas du tout le cas pour la sémantique ensembliste, lorsqu'on interprète toutes les formules atomiques par le même espace cohérent *fini* : dans le cas des entiers de Church, par exemple, une boîte coupée sur la formule conclusion du lien contraction d'arité arbitraire (aussi grande, en fait, que l'entier représenté) sera dupliquée un nombre de fois différent pour chaque entier. Dans ce cas on aura tendance à vouloir distinguer ces preuves entre elles. C'est une des raisons pour lesquelles nous avons privilégié la sémantique multiensembliste.

Il est tout de même intéressant, nous semble-t-il, de savoir que toutes ces identifications sont possibles, sans que pour autant la relation d'équivalence induite sur les preuves soit dégénérée. Ce qui nous conduit à formuler une conjecture. Elle n'est pas vraiment conforme à notre point de vue, et nous n'avons pas tenté de la démontrer.

Conjecture : Soit $\sim_0$ la relation d'équivalence sur les réseaux de *MELL* (pas forcément standard) définie par $R \sim_0 R'$ ssi $[R]_0 = [R']_0$, où R et R' sont des réseaux de même conclusion et $[R]_0$ est l'interprétation cohérente ensembliste de R obtenue en interprétant les atomes de R par l'espace cohérent de trame vide.

Soit $\sim$ une relation d'équivalence sur les réseaux de *MELL* (telle que si $R \sim R'$ alors R et R' ont le même séquent conclusion) satisfaisant les conditions (1) et (3) de la proposition 13.3.1, et telle que si $R \sim_0 R'$ alors $R \sim R'$. On a alors forcément $\sim = \sim_0$.

Pour en revenir à la question naïve ci-dessus, il nous semble que les réseaux de *MELL* ne sont pas en mesure de rendre compte de la condition d'uniformité que la sémantique cohérente multiensembliste parvient à capturer. Nous allons donc nous laisser conduire par cette dernière, dans la suite.

Concluons ce chapitre en remarquant que dans les deux contre-exemples à l'injectivité la présence des affaiblissements est essentielle, et que aucun des réseaux considérés n'est un réseau *polarisé*.

Chapitre 14

Fragments d'injectivité

On a montré dans le chapitre précédent que la conjecture 7.1.3 est fausse pour la sémantique cohérente, à cause de la propriété d'uniformité. En suivant l'idée que cette propriété doit bien avoir un sens syntaxique (et pourquoi pas "géométrique" ?), nous allons chercher des fragments de *MELL* pour lesquels la sémantique cohérente multiensembliste est injective.

Nous montrons que pour certains réseaux (standard, sans boîtes, sans affaiblissements, sans liens $\wp$), il existe une expérience injective (proposition 14.1.4). La démonstration suggère l'existence d'un lien entre la "connexité" (en un sens à préciser) du réseau et la propriété que nous voulons prouver. Le théorème précis qui découle de notre travail est le théorème 14.2.1.

La proposition 14.1.4 a aussi comme conséquence un théorème d'injectivité pour le fragment "faiblement polarisé" de LL défini dans 14.2.3 (le théorème 14.2.7). En particulier, ce dernier théorème permet de prouver l'injectivité de la sémantique cohérente multiensembliste pour le système *ILLU* de logique minimale (corollaire 14.2.8), c.à.d. pour le λ -calcul simplement typé.

Enfin, nous laissons présager d'une suite possible de notre analyse (la conjecture 14.2.10).

14.1 Le cas des contractions terminales

Dans tout le présent paragraphe, et sauf mention explicite du contraire, un réseau sera toujours un réseau standard, sans liens "!" (c.à.d. sans boîtes), sans liens $\wp$, et sans liens $?w$.

Nous allons montrer que si dans le réseau R tous les liens contraction sont terminaux, alors il existe une expérience injective de R (proposition 14.1.4). Nous voudrions que le lecteur prête une attention toute particulière au rôle

que joue la *connexité* de R dans notre argument : nous prétendons que c'est en fait cette propriété des preuves qui est à mettre en relation avec la propriété d'uniformité de la sémantique cohérente.

Si α est une arête de type atomique conclusion du lien axiome n du réseau R , nous noterons dans la suite $\alpha^\perp$ l'arête conclusion de n différente de α .

14.1.1. REMARQUE. Soient a et a' deux arêtes (différentes et) de même type du réseau R . Puisqu'il n'y a pas de liens affaiblissement dans R , le cas (3) du lemme 11.1.4 est à exclure. Par ailleurs, si $a \in G_{a'}$ ou bien si $a' \in G_a$, alors (puisque'il n'y a pas de boîtes dans R) forcément a ou a' est conclusion d'un lien *?co*. Le lemme en question nous dit alors que pour toute expérience injective e de R , $e(a) \neq e(a')$. Il n'y a donc que deux possibilités pour la relation de cohérence entre $e(a)$ et $e(a')$: $e(a) \sim e(a')$ ou bien $e(a) \frown e(a')$.

Une expérience injective d'un réseau R associe donc à une paire d'arêtes différentes et de même type un des deux éléments de l'ensemble $\{\frown, \sim\}$. Ceci justifie la notation $e(a, a')$ pour la valeur " $\frown$ " ou bien " $\sim$ " de e sur la paire d'arêtes de même type $\{a, a'\}$. Par définition d'expérience (6.2.1) la valeur de l'expérience injective e sur une paire d'arête $\{a, a'\}$ ne dépend que de la valeur que e prend sur les paires d'arêtes *atomiques* de même type (en fait "semblables", cf. définition 14.1.9) $\{\alpha, \alpha'\}$, où $\alpha \in G_a$ et $\alpha' \in G_{a'}$. En fait, la donnée d'une expérience injective d'un réseau R est *exactement* la donnée d'une fonction qui associe à toute paire d'arêtes de même type atomique un des deux éléments de l'ensemble $\{\frown, \sim\}$, la *seule* contrainte étant que si α et α' sont deux arêtes de même type atomique, alors la valeur de la fonction sur la paire $\{\alpha, \alpha'\}$ est différente de la valeur de la fonction sur la paire $\{\alpha^\perp, \alpha'^\perp\}$.

Rappelons enfin que la relation de cohérence est *binaire*, et on pourra donc définir une (prétendue) expérience e' à partir d'une expérience e en modifiant la relation de cohérence entre les arêtes de même type atomique α et α' , la seule conséquence étant que nous aurons aussi modifié la relation de cohérence entre les arêtes $\alpha^\perp$ et $\alpha'^\perp$.

14.1.2. REMARQUE. Nous parlerons souvent, dans la suite, d'un "couple d'arêtes", là où l'expression "paire d'arêtes" serait en fait plus appropriée. Ceci est dû au fait que la gestion des couples s'avère souvent plus simple que celle des paires.

14.1.3. REMARQUE. Soit R un réseau. A tout lien axiome l de conclusion α_l et $\alpha_l^\perp$ on associe un élément x_l de la trame de l'espace cohérent $\mathcal{A}$ de telle

sorte que si $l \neq l'$ alors $x_l \neq x_{l'}$.

La remarque 12.0.5 permet d'affirmer que pour que l'étiquetage des arêtes de R induit par la définition ci-dessus donne lieu à une expérience de R , *il faut et il suffit* que pour tout lien $?co$ de R de prémisses $a_1, \dots, a_k$ on ait $e(a_i) \sim e(a_j)$, et ce $\forall i, j \in \{1, \dots, k\}$.

Nous allons prouver la

14.1.4. PROPOSITION. *Si R est un réseau dont tous les liens $?co$ sont terminaux, alors il existe une expérience injective de R .*

Démonstration : C'est une conséquence de la proposition 14.1.5. □

La remarque 14.1.3 montre bien que pour prouver la proposition précédente, il suffira de prouver la proposition suivante. Le lecteur n'aura d'ailleurs aucun mal à se convaincre que les deux énoncés sont en fait équivalents.

14.1.5. PROPOSITION. *Soit R un réseau sans liens contraction. Il existe une expérience injective e de R telle que pour tout couple (a, a') d'arêtes conclusions de R de même type A , on a $e(a) \sim e(a')(\mathcal{A})$.*

Démonstration : C'est une conséquence de la proposition 14.1.18. □

Convention : *Dans tout le reste du paragraphe*, lorsque nous parlerons d'un réseau, il s'agira toujours d'un réseau ne contenant aucun lien $?co$. Autrement dit, il n'y aura plus, dans nos réseaux, que des liens axiome, tenseur, dérélction.

14.1.6. REMARQUE. L'absence de liens $?co$ dans un réseau R rend évidente l'existence d'une expérience injective de R . Ce qui n'est pas évident c'est que cette expérience satisfasse la conclusion de la proposition 14.1.5.

Nous allons commencer par nous familiariser avec ces objets que sont à présent nos réseaux. Dans la suite, nous passerons constamment des réseaux à leur "représentation duale" sous forme d'arbre, introduite dans la remarque suivante.

14.1.7. REMARQUE. (Représentation duale). Un réseau de conclusions les arêtes $a_1, \dots, a_n$ est constitué de n “blocs” (les graphes $G_{a_1}, \dots, G_{a_n}$), connectés entre eux par des liens axiome (et leurs conclusions). Ce cas ressemble au cas purement multiplicatif du chapitre 7, à la différence que dans nos réseaux il y a un unique graphe de correction (cf. définition 1.3.6) : le réseau lui-même.

Ceci suggère une autre représentation possible pour nos réseaux, une sorte de “représentation duale” : tout bloc est un noeud, et tout lien axiome de conclusions α et $\alpha^\perp$ est une arête connectant deux noeuds. Si R est un réseau, nous ferons référence à cette représentation de R comme à la “représentation duale” de R , que nous noterons $R^\star$.

Il est bien évident que pour tout réseau R , le graphe $R^\star$ est un arbre (c.à.d. qu’il est acyclique et connexe).

Si a et b sont deux arêtes conclusions de R , α et $\alpha^\perp$ sont deux arêtes de type atomique conclusions d’un même lien axiome telles que $\alpha \in G_a$ et $\alpha^\perp \in G_b$, alors nous noterons toujours a et b les noeuds de $R^\star$ correspondant aux blocs G_a et G_b , et nous noterons $\alpha\alpha^\perp$ (resp. $\overrightarrow{\alpha\alpha^\perp}$) l’arête non orientée (resp. orientée) de $R^\star$ connectant a et b .

Si Λ et Λ' sont deux chemins orientés de $R^\star$ t.q. le dernier noeud de Λ est le premier noeud de Λ' , alors on notera $\Lambda \star \Lambda'$ le chemin orienté de $R^\star$ dont le premier (resp. le dernier) noeud est le premier (resp. le dernier) noeud de Λ (resp. Λ').

14.1.8. DÉFINITION. (Chemins de $R^\star$). Si a et b sont deux arêtes conclusions de R , alors il existe un unique chemin non orienté (que nous noterons $\Theta_{a,b}$) de $R^\star$ connectant a et b . On notera $\overrightarrow{\Theta_{a,b}}$ (resp. $\overrightarrow{\Theta_{b,a}}$) le chemin orienté de $R^\star$ ayant comme premier noeud a (resp. b) et comme dernier noeud b (resp. a).

Que le lecteur nous pardonne pour la définition suivante : il s’agit de définir avec un peu de précision la notion (intuitivement évidente) de deux arêtes “pareilles” (mais *vraiment* pareilles) dans les graphes de deux arêtes de même type.

14.1.9. DÉFINITION. (arêtes semblables) Soit A une formule ne contenant que les connecteurs “?” et “ $\otimes$ ”, et soit A' une formule quelconque obtenue à partir de A en changeant le nom des variables propositionnelles de A de telle sorte que celles-ci soient deux à deux différentes dans A' .

Soient a_1 et a_2 deux arêtes de type A du réseau R . Soit G'_{a_1} l'arbre obtenu à partir de G_{a_1} en changeant le type des arêtes de G_{a_1} de telle sorte que maintenant l'arête a_1 soit de type A' . Remarquons d'une part que ceci est possible parce qu'il n'y a aucun lien $?co$ dans R , et d'autre part que G'_{a_1} est bien défini. Nous considérons ici que les arêtes de G'_{a_1} sont les mêmes que celles de G_{a_1} , mais avec des types différents. On définit G'_{a_2} de la même façon. Soit c_1 (resp. c_2) une arête de G_{a_1} (resp. de G_{a_2}). Nous dirons que c_1 et c_2 sont semblables lorsque, en tant qu'arêtes de G'_{a_1} et G'_{a_2} , c_1 et c_2 sont de même type.

14.1.10. REMARQUE. Soient a_1 et a_2 deux arêtes conclusions du réseau R . Si c_1 (resp. c_2) est une arête de G_{a_1} (resp. de G_{a_2}), et si c_1 et c_2 sont semblables, alors a_1 et a_2 le sont aussi (c.à.d. que a_1 et a_2 sont deux conclusions de même type).

14.1.11. REMARQUE. La définition de la relation de cohérence dans les espaces $\mathcal{A} \otimes \mathcal{B}$ et $?B$ permet d'affirmer que pour tout réseau R et pour toute expérience injective e de R , on a :

pour tout couple a, a' de conclusions de même type de R , $e(a) \sim e(a')$ ssi il existe un couple $(\alpha, \alpha') \in G_a \times G_{a'}$ d'arêtes semblables de type atomique t.q. $e(\alpha) \sim e(\alpha')$.

Nous introduisons, dans la définition suivante, la notion de (C) -couple : deux arêtes atomiques semblables de R vont constituer un (C) -couple pour un couple de conclusions de R lorsqu'elles “participeront” à la connexion entre ces deux conclusions (dans R^*).

14.1.12. DÉFINITION. (couple de connexité) Soit R un réseau et a et a' deux arêtes de même type conclusions de R .

On dira que le couple d'arêtes semblables de type atomique $(\alpha, \alpha') \in G_a \times G_{a'}$ est un couple de connexité (ou plus succinctement un C -couple) pour (a, a') lorsque le chemin $\Theta_{a,a'}$ de R^* connectant a et a' contient l'arête $\alpha\alpha^\perp$ ou/et l'arête $\alpha'\alpha'^\perp$.

14.1.13. REMARQUE. Pour tout couple (a, a') d'arêtes conclusions de même type d'un réseau R , il existe toujours un (C) -couple pour (a, a') , et il en existe au plus deux.

Remarquons que lorsque le type de a et de a' ne contient aucune occurrence du connecteur $\otimes$, il n'y aura qu'un seul (C) -couple pour (a, a') .

Remarquons aussi que si (α, α') est l'unique (C) -couple pour (a, a') , alors $(\alpha^\perp, \alpha'^\perp)$ ne peut pas être un (C) -couple. Cette remarque est importante : nous en ferons usage dans la démonstration de la proposition 14.1.18.

14.1.14. DÉFINITION. (couple de chemins semblables) Soit R un réseau et a, a' deux arêtes conclusions de R et de même type. Soit n un entier positif non nul. Soit Φ (resp. Φ') un chemin orienté de R^* issu de a (resp. de a') dont les arêtes sont $\alpha_1 \alpha_1^\perp, \dots, \alpha_n \alpha_n^\perp$ (resp. $\alpha'_1 \alpha'_1{}^\perp, \dots, \alpha'_n \alpha'_n{}^\perp$). Nous dirons que (Φ, Φ') est un couple de chemins semblables issu de (a, a') lorsque :

- les arêtes α_i et α'_i sont semblables $\forall i \in \{1, \dots, n\}$
- les arêtes $\alpha_i^\perp$ et $\alpha'_i{}^\perp$ sont semblables $\forall i \in \{1, \dots, n-1\}$.

Nous dirons que le couple de chemins semblables (Φ, Φ') est maximal, lorsque les arêtes $\alpha_n^\perp$ et $\alpha'_n{}^\perp$ ne sont pas semblables.

La proposition suivante est une conséquence de la connexité de l'arbre R^* , et c'est l'ingrédient essentiel dans la preuve de la proposition 14.1.18. Pour en donner une preuve précise nous associons, à tout chemin orienté de R^* , un entier sur lequel on pourra raisonner par induction.

14.1.15. DÉFINITION. Soient c et d deux arêtes conclusions du réseau R , et soit $\gamma\gamma^\perp$ l'arête de $\Theta_{c,d}$ dont d est une des extrémités.

Commençons par remarquer qu'on peut facilement ordonner les sous-arbres de R^* contenant d et ne contenant pas l'arête $\gamma\gamma^\perp$: par exemple en comptant le nombre d'arêtes de ces sous-arbres. Il est bien clair que l'ensemble de ces sous-arbres ainsi ordonné a un élément maximal : notons-le A_d .

On peut alors associer un entier au chemin orienté $\Theta_{c,d}$: le nombre d'arêtes de l'arbre A_d , que l'on notera $\|\overrightarrow{\Theta_{c,d}}\|$.

Remarquons que pour deux arêtes c et d comme ci-dessus, on a $\|\overrightarrow{\Theta_{c,d}}\| = 0$ ssi d est une feuille de R^* , c.à.d. ssi le type de d ne contient pas d'occurrences du connecteur $\otimes$.

14.1.16. PROPOSITION. (couple de chemins semblables maximal). Soient a et a' deux arêtes de même type conclusions du réseau R , et supposons qu'il existe pour (a, a') deux (C) -couples.

Il existe un couple de chemins semblables maximal (Φ, Φ') issu de (a, a') et t.q. Φ est un sous-chemin orienté du chemin orienté $\overrightarrow{\Theta_{a,a'}} \star \Phi'$.

Démonstration : Soit (α, α') un (C) -couple pour (a, a') et supposons que $\overrightarrow{\alpha\alpha^\perp}$ soit la première arête de $\overrightarrow{\Theta_{a,a'}}$. Appelons alors $\overrightarrow{\alpha_n\alpha_n^\perp}$ la dernière arête de $\overrightarrow{\Theta_{a,a'}}$: puisqu'il y a deux (C) -couples pour (a, a') , on a $\alpha' \neq \alpha_n^\perp$.

Soit b (resp. b') l'arête conclusion de R t.q. $\alpha^\perp \in G_b$ (resp. $\alpha'^\perp \in G_{b'}$).

Si $\alpha^\perp$ et $\alpha'^\perp$ ne sont pas semblables, alors la proposition est vraie. Si elles le sont, alors b et b' sont deux conclusions de même type de R (par la remarque 14.1.10). Puisque $\alpha' \neq \alpha_n^\perp$, le noeud b' n'est pas un noeud de $\overrightarrow{\Theta_{a,a'}}$ et l'arête $\overrightarrow{\alpha'\alpha'^\perp}$ n'est pas une arête de $\overrightarrow{\Theta_{a,a'}}$. Plus précisément, on a $\overrightarrow{\Theta_{a,a'}} \star \overrightarrow{\alpha'\alpha'^\perp} = \overrightarrow{\alpha\alpha^\perp} \star \overrightarrow{\Theta_{b,b'}}$.

Soit $\overrightarrow{\beta\beta^\perp}$ la première arête de $\overrightarrow{\Theta_{b,b'}}$: on a $\beta \neq \alpha^\perp$. Soit $\beta' \in G_{b'}$ l'arête semblable à β : on a $\beta' \neq \alpha'^\perp$. Le couple (β, β') est un (C) -couple pour (b, b') , et il existe deux (C) -couples pour (b, b') : (β, β') et $(\alpha^\perp, \alpha'^\perp)$. On est donc revenu à la situation de départ, mais cette fois avec le couple (b, b') . La raison pour laquelle il va forcément falloir s'arrêter un jour, c'est que $R^\star$ est un arbre.

Plus formellement, on procède par induction sur $\|\overrightarrow{\Theta_{a,a'}}\|$.

Si $\|\overrightarrow{\Theta_{a,a'}}\| = 0$, alors a' est une feuille de l'arbre $R^\star$. Le type de a' (qui est aussi celui de a) est alors atomique, éventuellement prefixé par un certain nombre de "?". De toutes façons a sera aussi une feuille de $R^\star$, et l'hypothèse de la proposition est alors fausse, puisque dans ce cas il existe un unique (C) -couple pour (a, a') .

Si $\|\overrightarrow{\Theta_{a,a'}}\| > 0$, alors on vient de voir que les hypothèses de la proposition sont satisfaites pour les noeuds b et b' de $R^\star$. On aura par ailleurs $\|\overrightarrow{\Theta_{a,a'}}\| > \|\overrightarrow{\Theta_{b,b'}}\|$, ce qui nous permettra d'appliquer l'hypothèse de récurrence : appelons (Ψ, Ψ') le couple de chemins semblables maximal issu de (b, b') tel que Ψ est un sous-chemin de $\overrightarrow{\Theta_{b,b'}} \star \Psi'$. Appelons Φ (resp. Φ') le chemin $\overrightarrow{\alpha\alpha^\perp} \star \Psi$ (resp. $\overrightarrow{\alpha'\alpha'^\perp} \star \Psi'$). Le couple (Φ, Φ') est bien un couple de chemins semblables maximal issu de (a, a') . Par ailleurs $\overrightarrow{\Phi} = \overrightarrow{\alpha\alpha^\perp} \star \overrightarrow{\Psi}$ est un sous-chemin de $\overrightarrow{\alpha\alpha^\perp} \star \overrightarrow{\Theta_{b,b'}} \star \overrightarrow{\Psi'} = \overrightarrow{\Theta_{a,a'}} \star \overrightarrow{\alpha'\alpha'^\perp} \star \overrightarrow{\Psi'} = \overrightarrow{\Theta_{a,a'}} \star \overrightarrow{\Phi'}$. $\square$

La remarque suivante va nous être utile dans la démonstration de la proposition qui la suit.

14.1.17. REMARQUE. Reprenons les notations de la proposition précédente. Soient c et c' deux conclusions de même type de R . Si $(\gamma, \gamma') \in G_c \times G_{c'}$ est un couple d'arêtes semblables de type atomique t.q. $\xrightarrow{\gamma} \gamma\gamma^\perp$ (resp. $\xrightarrow{\gamma'} \gamma'\gamma'^\perp$) est une arête de Φ (resp. Φ'), alors (γ, γ') est un (C) -couple pour (c, c') .

La proposition suivante, qui aura comme conséquence la proposition 14.1.5 (et donc la proposition 14.1.4 que nous cherchons à prouver) met bien l'accent sur l'importance de la connexité. Elle prouve que la solution à notre problème pour un couple de conclusions de même type (a, a') est fournie *précisément* par les couples d'arêtes semblables qui assurent la connexion entre a et a' : les (C) -couples.

14.1.18. PROPOSITION. Soit R un réseau. Il existe une expérience injective e de R telle que :
pour tout couple (a, a') d'arêtes conclusions de R de même type, il existe un C -couple (α, α') pour (a, a') satisfaisant $e(\alpha) \sim e(\alpha')$.

Démonstration : Soit k le nombre de paires d'arêtes conclusions de même type du réseau R . Soit e une expérience injective quelconque de R et k_e le nombre des paires d'arêtes $\{c, c'\}$ conclusions de même type de R pour lesquelles il existe un (C) -couple $(\alpha_c, \alpha_{c'}) \in G_c \times G_{c'}$ t.q. $e(\alpha_c) \sim e(\alpha_{c'})$.

On va prouver, par induction sur $k - k_e$, qu'il existe une expérience injective e' de R satisfaisant la conclusion de la proposition.

Si $k - k_e = 0$, alors e est l'expérience injective de R cherchée.

Autrement, il existe un couple (a, a') de conclusions de même type de R t.q. pour le ou les (C) -couple(s) $(\alpha_a, \alpha_{a'})$ pour (a, a') , on a $e(\alpha_a) \sim e(\alpha_{a'})$. On va construire, à partir de e , une expérience injective e' de R t.q. $k_{e'} > k_e$. Il suffira ensuite d'appliquer l'hypothèse d'induction.

Fixons un (C) -couple (α, α') pour (a, a') . Si (α, α') est le seul (C) -couple pour (a, a') , alors (voir remarque 14.1.13) $(\alpha^\perp, \alpha'^\perp)$ ne peut pas être un (C) -couple. Dans ce cas, il suffira de définir l'expérience e' comme l'expérience e sauf sur $\{\alpha, \alpha'\}$ et sur $\{\alpha^\perp, \alpha'^\perp\}$: on posera dans ce cas $e'(\alpha) \sim e'(\alpha')$ (et donc $e'(\alpha^\perp) \sim e'(\alpha'^\perp)$). On aura alors $k_{e'} = k_e + 1$.

Si par contre (α, α') n'est pas le seul (C) -couple pour (a, a') , alors on peut appliquer la proposition 14.1.16 : soit (Φ, Φ') le couple de chemins semblables maximal t.q. $\xrightarrow{\Phi} \Phi$ est un sous-chemin de $\Theta_{a, a'} \star \Phi'$.

Supposons que Φ et Φ' contiennent $k + 1$ arêtes différentes ($k \geq 0$). Appelons $(b_1, b'_1), \dots, (b_k, b'_k)$ les couples d'arêtes conclusions de R et $(\beta_1, \beta'_1), \dots, (\beta_k, \beta'_k)$ les couples d'arêtes atomiques de R t.q. :

- les noeuds traversés par Φ' sont, dans l'ordre, $a', b'_1, \dots, b'_k, b'_{k+1}$
- les noeuds traversés par Φ sont, dans l'ordre, $a, b_1, \dots, b_k, b_{k+1}$
- $\beta'_i \in G_{b'_i}$ (resp. $\beta_i \in G_{b_i}$) et $\beta'^{\perp}_i \in G_{b'_{i+1}}$ (resp. $\beta^{\perp}_i \in G_{b_{i+1}}$), pour $i \in \{1, \dots, k\}$
- $\forall i \in \{1, \dots, k\}$ (resp. $\forall i \in \{1, \dots, k-1\}$), β'_i et β_i (resp. $\beta'^{\perp}_i$ et $\beta^{\perp}_i$) sont semblables, tandis que $\beta'^{\perp}_k$ et $\beta^{\perp}_k$ ne sont pas semblables
- $\overrightarrow{\alpha\alpha^{\perp}}$ (resp. $\overrightarrow{\alpha'\alpha'^{\perp}}$) est l'arête de Φ' (resp. de Φ) connectant a à b_1 (resp. a' à b'_1)
- $\forall i \in \{1, \dots, k\}$, $\overrightarrow{\beta'_i\beta'^{\perp}_i}$ (resp. $\overrightarrow{\beta_i\beta^{\perp}_i}$) est l'arête de Φ' (resp. de Φ) connectant b'_i et b'_{i+1} (resp. b_i et b_{i+1}).

On définit alors e' sur toute paire d'arêtes atomiques de même type $\{\delta, \delta'\}$ de R de la façon suivante :

- si $\{\delta, \delta'\} \notin \{\{\alpha, \alpha'\}, \{\alpha^{\perp}, \alpha'^{\perp}\}\} \cup \{\{\beta_i, \beta'_i\}, \{\beta^{\perp}_i, \beta'^{\perp}_i\} : i \in \{1, \dots, k\}\}$, alors $e'(\delta, \delta') = e(\delta, \delta')$
- $e'(\alpha) \sim e'(\alpha')$ (donc $e'(\alpha^{\perp}) \sim e'(\alpha'^{\perp})$)
- $\forall i \in \{1, \dots, k\}$, $e'(\beta_i) \sim e'(\beta'_i)$ (donc $e'(\beta^{\perp}_i) \sim e'(\beta'^{\perp}_i)$).

Montrons maintenant que $k_{e'} > k_e$.

Si $\{c, c'\} \notin \{\{a, a'\}, \{b_1, b'_1\}, \dots, \{b_k, b'_k\}\}$, alors $\forall \delta \in G_c$ et $\forall \delta' \in G_{c'}$ avec δ et δ' semblables de type atomique, on a $e'(\delta, \delta') = e(\delta, \delta')$ (n'oublions pas que $\beta'^{\perp}_k \in G_{b'_{k+1}}$ et $\beta^{\perp}_k \in G_{b_{k+1}}$ ne sont pas semblables). Ce qui signifie que si (δ, δ') est un (C) -couple pour (c, c') t.q. $e(\delta) \sim e(\delta')$, ce sera toujours le cas pour e' .

Si $\{c, c'\} \in \{\{b_1, b'_1\}, \dots, \{b_k, b'_k\}\}$, alors par la remarque 14.1.17 on sait que $\forall i \in \{1, \dots, k\}$ le couple (β_i, β'_i) est un (C) -couple pour (b_i, b'_i) , et par définition de e' on a $e'(\beta_i) \sim e'(\beta'_i)$.

Si $(c, c') = (a, a')$, alors on a par définition de e' : $e'(\alpha) \sim e'(\alpha')$ (où (α, α') est un (C) -couple).

Dans tous les cas, il est donc bien clair que $k_{e'} \geq k_e + 1$. $\square$

14.2 Conclusions

Nous allons maintenant énoncer le théorème que notre méthode nous a amené à prouver (le théorème 14.2.1). Il prouve l'injectivité pour tous les sous-systèmes de *MELL* qui satisfont une certaine condition.

Nous montrons ensuite qu'une simple variante de la proposition 14.1.4 (et une remarque amusante) permet(tent) de prouver l'injectivité de la sémantique cohérente multiensembliste pour un fragment significatif de *MELL* (théorème 14.2.7). En particulier, le système *ILLU* de logique minimale (c.à.d. le λ -calcul

simplement typé) fait partie de notre fragment, et la sémantique cohérente multiensembliste est donc injective pour ILL .

Nous concluons en remarquant que si le théorème 14.2.7 est une façon d'exprimer notre résultat, ce n'est certainement pas la seule : c'est ce que, techniquement, nous sommes capables de démontrer à ce jour, mais il est vraisemblable que la méthode puisse s'étendre à des cadres bien plus amples, comme nous le conjecturons dans 14.2.10.

Nous oublierons, dans ce dernier paragraphe, les conventions sur les réseaux adoptées jusqu'à présent dans ce chapitre.

Soit R un réseau de $MELL$. On notera R_0 le réseau standard associé à sa forme normale (définition 7.1.1), $L(R_0)$ le linéarisé de R_0 (définition 11.2.1), et enfin (lorsque $L(R_0)$ ne contient pas d'affaiblissements) on notera $L(R_0)^\wp$ l'ensemble des réseaux obtenus à partir de $L(R_0)$ en éliminant les liens $\wp$ (remarque 12.0.9).

Dans les énoncés suivants, nous utilisons les notations de la définition 7.1.2. Pour un réseau R , on notera $[R]_{coh m}$ la sémantique cohérente multiensembliste de R .

14.2.1. THÉORÈME. Soit R un réseau de $MELL$ tel que R_0 ne contient pas de liens affaiblissement et il existe un élément de $L(R_0)^\wp$ dont tous les liens contraction sont terminaux.

Si R' est un réseau de même conclusion que R et si $[R]_{coh m} = [R']_{coh m}$, alors $R \simeq_{\beta\eta} R'$.

Démonstration : Soit $L(R_0)^-$ l'élément de $L(R_0)^\wp$ ne contenant que des liens contraction terminaux. Par la proposition 14.1.4, il existe une expérience injective de $L(R_0)^-$. Par la proposition 12.0.10, il existe une expérience injective de $L(R_0)$. Soit alors $n > \sup(1, h)$ où h est l'arité maximale des liens contraction de R_0 . Par la proposition 11.2.4, il existe une n -expérience obsessionnelle injective de R_0 . Soit R'_0 le réseau standard associé à R' . Par le théorème 10.2.4, on a $R_0 = R'_0$, c.à.d. $R \simeq_{\beta\eta} R'$. $\square$

14.2.2. THÉORÈME. Soit R un réseau de $MELL$ tel que R_0 ne contient pas de liens affaiblissement. Supposons qu'il existe un élément $L(R_0)^-$ de $L(R_0)^\wp$ pour lequel il existe une expérience injective.

Si R' est un réseau de même conclusion que R et si $[R]_{coh m} = [R']_{coh m}$, alors $R \simeq_{\beta\eta} R'$.

Démonstration : Elle est identique à celle du théorème précédent, à la différence que nous supposons ici l'existence de l'expérience injective de $L(R_0)^-$, alors que nous l'avons prouvée ci-dessus. $\square$

Nous allons définir la notion de formule “faiblement polarisée” qui est à mettre en relation avec la définition 15.2.3 de formule polarisée et la définition 15.5.1 de formule fortement polarisée.

14.2.3. DÉFINITION. (Formules faiblement polarisées) Une formule propositionnelle P (resp. N) de LL est dite faiblement positive (resp. faiblement négative) lorsqu'elle est construite de la façon suivante (X étant une formule atomique) :

$$\begin{array}{lcl} P ::= & X & | \quad P \otimes P \quad | \quad P \otimes !N \quad | \quad !N \otimes P \\ N ::= & X & | \quad N \wp N \quad | \quad ?P \wp N \quad | \quad N \wp ?P \end{array}$$

Nous dirons qu'une formule est faiblement polarisée lorsqu'elle est faiblement positive ou bien faiblement négative.

On dira qu'un réseau de *MELL* est faiblement polarisé lorsque les types de ses conclusions sont tous des sous-formules de formules faiblement polarisées.

14.2.4. REMARQUE. La particularité de la définition ci-dessus par rapport aux deux autres citées est qu'on ne suppose rien sur les atomes, qui sont tous à la fois faiblement positifs et faiblement négatifs.

Ceci signifie, en particulier, qu'une formule faiblement polarisée A n'est pas équivalente à $?A$ ni à $!A$, et ce contrairement aux formules polarisées (voir la proposition 15.2.4).

14.2.5. LEMME. Soit R un réseau standard de *MELL* sans boîtes. Si R satisfait le propriété :

(P) toute arête de type $?F$ (F formule de LL) est (modulo un lien $?co$) prémisses d'un lien $\wp$ ou bien terminale, et tout lien $\wp$ a au moins une prémisses qui n'est pas de type $?F$ (F formule de LL)
alors il existe une expérience injective de R .

Démonstration : Par induction sur le nombre k de liens $\wp$ de R .

Si $k = 0$, alors toute arête de type $?F$ de R est terminale (modulo un lien contraction). Soit R' le réseau obtenu à partir de R en effaçant tous les liens affaiblissements (nécessairement terminaux) de R . Tout lien contraction de R' est terminal, et il n'y a pas de lien $\wp$ dans R' : par la proposition 14.1.4,

il existe une expérience injective de R' . Et donc une expérience injective de R .

Si $k > 0$, soit alors a une conclusion de R t.q. il existe un lien $\wp$ dans G_a . Soit n un lien $\wp$ de G_a qui est un “parmi les plus proches de a ”. Cette expression a clairement un sens (il est facile de mesurer la distance à a dans l'arbre G_a). On peut appliquer à *ce lien* $\wp$, la procédure d'élimination décrite au chapitre 12, et ce malgré l'éventuelle présence d'affaiblissements. Nous allons le faire judicieusement, de telle sorte que le réseau R' ainsi obtenu satisfasse encore la propriété (P) . Si jamais (avec les notations de la définition 12.0.7) une des deux prémisses des a_i (par exemple c_i) est de type $C = ?E$, alors on va appliquer la procédure de telle sorte à obtenir un réseau de conclusion $\Gamma \setminus A, A[D/C \wp D], ?E$ (et pas $\Gamma \setminus A, A[C/C \wp D], ?D$). Il est capital de remarquer que c'est parce que (par hypothèse) $D \neq ?F$ pour toute formule F , que la propriété (P) sera toujours satisfaite par R' . La conclusion s'obtient alors en appliquant l'hypothèse d'induction. $\square$

14.2.6. PROPOSITION. Soit R un réseau faiblement polarisé. Il existe toujours une expérience injective de $L(R_0)$.

Démonstration : Il suffit de remarquer que $L(R_0)$ satisfait la propriété (P) du lemme 14.2.5, ce qui est évident. $\square$

14.2.7. THÉORÈME. Soit R un réseau faiblement polarisé. Si R' est un réseau de même conclusion que R et si $[R]_{coh m} = [R']_{coh m}$, alors $R \simeq_{\beta\eta} R'$.

Démonstration : La proposition 14.2.6 permet d'affirmer qu'il existe une expérience injective de $L(R_0)$. Soit alors $n > \sup(1, h)$ où h est l'arité maximale des liens contraction de R_0 . Par la proposition 11.2.4, il existe une n -expérience obsessionnelle injective de R_0 . Soit R'_0 le réseau standard associé à R' . Par le théorème 10.2.4, on a $SPL(R_0) = SPL(R'_0)$.

Remarquons alors que puisque R_0 et R'_0 sont faiblement polarisés, tout lien $?w$ de R_0 ou R'_0 qui n'est pas terminal est prémisses d'un lien $\wp$; et tout lien $\wp$ de ces deux réseaux a au moins une prémisses qui n'est pas conclusion d'un lien $?w$.

La remarque “amusante” dont il était question au début de ce paragraphe, est la suivante : on peut affirmer que, *dans ce cas*, même en présence d'affaiblissements, la caractérisation des boîtes que fournit la preuve de la proposition 10.0.7 est toujours valable pour R_0 et R'_0 .

On en déduira que $R_0 = R'_0$, c.à.d. $R \simeq_{\beta\eta} R'$. $\square$

14.2.8. COROLLAIRE. La sémantique cohérente multiensembliste est injective pour le fragment intuitionniste $ILLU$ de la logique unifiée LU de Girard ([Girard 93]).

Démonstration : Le système $ILLU$ est en fait le fragment t de LK^η (défini dans [DJS 97] et dont il sera question dans le chapitre 15), et Danos, Joinet et Schellinx ont prouvé que pour ce fragment la traduction de Girard $A \rightarrow B = !A \multimap B$ induit une sémantique dénotationnelle pour $ILLU$. Il suffit alors de remarquer que cette traduction n'utilise que des formules faiblement polarisées, et appliquer ensuite le théorème 14.2.7. $\square$

14.2.9. REMARQUE. On aurait pu donner des preuves “moins constructives” de la proposition 14.1.18, mais il nous semble que la preuve que nous avons donnée est plus intéressante, puisqu'elle suggère que la propriété de connexité des réseaux est liée à la propriété d'uniformité de la sémantique cohérente.

Un autre intérêt de cette preuve est qu'on peut espérer l'étendre au cas des réseaux fortement polarisés sans constantes (voir définition 15.5.1). Ce résultat serait particulièrement intéressant car il permettrait de prouver l'injectivité de la P -sémantique (définie dans le chapitre 15) pour le fragment multiplicatif du système LC de Girard (voir définition 15.5.3), comme nous le précisons dans la remarque 15.5.4.

Nous allons terminer en esquisant une preuve d'injectivité de la sémantique cohérente multiensembliste pour les réseaux fortement polarisés, qui seront toujours, dans les lignes qui suivent, sans additifs, sans affaiblissements, et sans constantes.

On procède exactement comme dans le cas faiblement polarisé jusqu'à l'étape d'élimination des liens $\wp$. La présence de formules telles que $?X \wp ?X$ parmi les formules fortement polarisées, nous empêchera d'aboutir à un réseau dont les liens $?co$ sont tous terminaux.

Si R est un réseau fortement polarisé, le lecteur se rendra compte facilement qu'on peut obtenir, en éliminant (judicieusement) les $\wp$ de $L(R_0)$ et en effaçant les liens $?co$ terminaux, un réseau R^- qui sera toujours constitué d'un certain nombre de “blocs” (comme dans la remarque 14.1.7). La différence essentielle avec le cas faiblement polarisé étant que les connexions entre ces blocs s'effectueront soit par un lien axiome et ses conclusions (comme dans le cas que nous avons analysé) soit par un “faisceau” de liens axiomes. Soyons un peu plus précis : “en haut” de chaque bloc nous n'aurons

plus forcément des arêtes de type atomique, mais aussi des conclusions de liens $?co$ dont les prémisses sont de type $?X$ (avec X atomique).

La représentation duale de la remarque 14.1.7 s'étend donc très bien au cas des réseaux fortement polarisés, à la différence que le graphe $(R^-)^*$ n'est plus un arbre mais... un réseau !

La représentation duale $(R^-)^*$ du réseau R^- sera un graphe apparié satisfaisant la condition (ACC) : pour tout choix S d'une prémisses de chaque lien $?co$ on obtient un graphe $S((R^-)^*)$ qui est un arbre.

La même technique de démonstration semble s'étendre au cas où $(R^-)^*$ satisfait (ACC) , mais nous n'avons pas encore pu le prouver.

Il est donc très naturel de conclure ce chapitre par la conjecture qui suit, et dont la conséquence immédiate (à la lumière du théorème 14.2.2) serait l'injectivité de la sémantique cohérente multiensembliste pour les réseaux fortement polarisés, et donc pour le fragment multiplicatif (et sans affaiblissements) de LC (par la remarque 15.5.4).

14.2.10 Conjecture

Soit R un réseau fortement polarisé et soit R^- le réseau (dont il était question ci-dessus) obtenu à partir de $L(R_0)$ en éliminant judicieusement les liens $\wp$ et en effaçant les liens $?co$ terminaux.

Il existe pour R^- une expérience injective telle que pour tout couple (a, a') d'arêtes conclusions de R^- de même type A , $e(a) \sim_e(a')(\mathcal{A})$.

Chapitre 15

Logique linéaire polarisée et logique classique

La notion de polarité fut introduite par J.-Y. Girard dans [Girard 91b]. La notion de focalisation fut clairement cernée pour la première fois dans [AndrPar 91], mais nous en avons pris connaissance à travers les travaux sur la logique classique [Girard 91b] (“le bénitier” de LC) et [DJS 97] (“la η -contrainte” de LK^η).

Les raisons pour considérer le fragment polarisé et focalisé de LL sont légions. La première nous vient sans aucun doute de [Girard 91b] : ce fragment (qui est un *tout petit* fragment de LL) est suffisant pour coder la logique classique, *et en plus* pour définir une procédure d’élimination des coupures munie d’une sémantique dénotationnelle pour la logique classique. Puisque la logique classique est tout de même LA logique, à elle toute seule cette motivation justifierait déjà amplement l’étude du fragment en question (en tous cas aux yeux de l’auteur de cette thèse).

L’étude de Girard a été étendue (tout au moins sémantiquement) au second ordre dans [Quatrini 95]. Plus récemment, O. Laurent a montré dans [Laurent 99] que dans le fragment polarisé la séquentialisation des réseaux avec poids de Girard ([Girard 95a]) est beaucoup plus simple. Nous montrerons dans [LaurQuatrTorto 99] que la normalisation est aussi beaucoup plus simple dans ce fragment, et ce indépendamment de la syntaxe choisie pour représenter les additifs (réseaux avec poids de Girard, multiboîtes du chapitre 5, ou simplement boîtes additives du chapitre 1). Voir aussi à ce sujet la remarque 15.2.8.

Le contenu de ce chapitre est le fruit d’une collaboration avec Myriam Quatrini (commencée avec [QuatrTorto 96]), et d’une collaboration plus tardive

avec Olivier Laurent, le tout étant contenu dans [LaurQuatrTorto 99]. Nous ferons aussi référence à une collaboration avec J.-B. Joinet et H. Schellinx, dont le contenu est exposé dans [JST 95] et [JST 98].

Les trois premiers paragraphes contiennent une présentation, assez succincte mais exhaustive, du cadre dans lequel nous exposerons les résultats des deux derniers paragraphes.

Dans le paragraphe 15.4, nous définissons une contrainte sur les preuves classiques (la ρ -contrainte) qui, avec la η -contrainte de [DJS 97], fournit un système $(LK_{pol}^{\eta,\rho})$ complet pour la prouvabilité classique (théorème 15.4.9) et stable par élimination des coupures (théorème 15.4.11). Un des intérêts d'un tel système est que le P -plongement de Girard est une *décoration* de $LK_{pol}^{\eta,\rho}$, dont l'image est précisément la logique linéaire polarisée et focalisée (proposition 15.4.12). Ceci permet de montrer que le P -plongement induit une sémantique dénotationnelle de la logique classique (théorème 15.4.14), grâce à laquelle une étude plus syntaxique des isomorphismes ayant présidé à la définition du système LC de Girard est possible au paragraphe 15.5. C'est toujours dans le système $LK_{pol}^{\eta,\rho}$ (plus précisément dans son fragment LC^ρ) que la question de l'injectivité de la sémantique pour la logique classique semble, à la lumière des résultats exposés dans les chapitres précédents, pouvoir trouver une réponse positive (remarque 15.5.4).

Nous utiliserons, dans la suite, la terminologie usuelle pour les règles du calcul des séquents : nous supposerons que le lecteur est familier avec la notion de conclusion et prémisses(s) d'une règle, ainsi qu'avec celle de formule active dans une règle.

Lorsque nous parlerons de "sémantique d'une preuve π ", celle-ci étant écrite dans un calcul des séquents (de LL) quel qu'il soit, nous nous référerons en fait à la sémantique du réseau associé à π , suivant la définition du chapitre 6. Dans ce chapitre, il s'agira toujours (sauf mention explicite du contraire) d'une quelconque des trois sémantiques présentées au chapitre 6 : la sémantique cohérente ensembliste, la sémantique cohérente multiensembliste, et la sémantique relationnelle (multiensembliste).

15.1 Le calcul des séquents classique LK

Nous présentons ci-dessous une version monolatère du calcul des séquents de Gentzen, qui tient compte du *style* multiplicatif ou additif des règles. C'est bien sûr une notion qui nous vient de LL, et qui n'a pas de véritable importance tant que l'on s'intéresse à des questions de *prouvabilité*. Mais nous verrons que le style d'une règle n'est point anodin, à partir du moment

où les objets sur lesquels on travaille sont les *preuves* et leur normalisation (cf. 15.3.3).

On notera $\vee_m$ (resp. $\wedge_m$) la disjonction (resp. la conjonction) multiplicative, et $\vee_a$ (resp. $\wedge_a$) la disjonction (resp. la conjonction) additive.

Nous avons aussi deux constantes, que nous noterons V et $\neg V$ (et qui correspondent aux constantes 1 et $\perp$ de LL). Nous allons nous restreindre au fragment propositionnel de LK , auquel nous rajoutons les règles pour les deux constantes V et $\neg V$, pour les raisons qui sont exposées au début du paragraphe 15.2. Nous noterons dans la suite LK ce calcul des séquents.

Le fragment propositionnel de LK

Axiome et coupure :

$$(Ax) \quad \vdash A, \neg A \qquad (cut) \quad \frac{\vdash \Gamma, A \quad \vdash \neg A, \Delta}{\vdash \Gamma, \Delta}$$

Règles logiques multiplicatives :

$$(\wedge_m) \quad \frac{\vdash \Gamma, A \quad \vdash \Delta, B}{\vdash \Gamma, \Delta, A \wedge_m B} \qquad (\vee_m) \quad \frac{\vdash \Gamma, A, B}{\vdash \Gamma, A \vee_m B}$$

Règles logiques additives :

$$(\vee_a) \quad \frac{\vdash \Gamma, A}{\vdash \Gamma, A \vee_a B} \quad \frac{\vdash \Gamma, B}{\vdash \Gamma, A \vee_a B} \qquad (\wedge_a) \quad \frac{\vdash \Gamma A \quad \vdash \Gamma, B}{\vdash \Gamma, A \wedge_a B}$$

Règles structurelles :

$$(W) \quad \frac{\vdash \Gamma}{\vdash \Gamma, A} \qquad (C) \quad \frac{\vdash \Gamma, A, A}{\vdash \Gamma, A}$$

Axiomes pour les constantes :

$$(V) \quad \vdash V \qquad (\neg V) \quad \frac{\vdash \Gamma}{\vdash \Gamma, \neg V}$$

15.1.1. DÉFINITION. (Formules, règles, classiques (ir)réversibles).

Les règles réversibles (resp. irréversibles) de LK sont les règles introduisant les connecteurs $\vee_m, \wedge_a$ (resp. $\wedge_m, \vee_a$), et la règle introduisant la constante $\neg V$ (resp. V).

Toute formule qui est la conclusion d'une règle réversible (resp. irréversible) est dite réversible (resp. irréversible).

15.1.2. REMARQUE. La notion de formule réversible (et irréversible) est bien connue en logique, une définition possible étant “la formule A est réversible ssi pour toute preuve π du séquent Γ, A , il existe une preuve π^r de Γ, A dont la dernière règle est une règle de conclusion A ”. Nous verrons dans la suite que la notion de réversibilité est en fait beaucoup plus forte que celle que nous venons de donner : si on s’y prend bien, non seulement la preuve π^r existe, mais elle a aussi la *même sémantique* que π (proposition 15.4.4).

15.1.3. DÉFINITION. (Formules classiques principales). Une (occurrence de) formule dans une preuve π de LK est dite principale lorsque c’est la conclusion d’une règle logique ($\forall_m, \wedge_m, \vee_a, \wedge_a$) ou bien la conclusion d’un axiome.

Lorsque nous souhaiterons être plus précis, nous dirons que la formule A est principale dans une règle logique, ou bien que A est principale dans un axiome.

15.2 La logique linéaire polarisée et focalisée LL_p^f

La logique linéaire polarisée et focalisée (LL_p^f) est essentiellement unique, en ce sens que les propriétés qui président à sa définition sont claires : polarisation et focalisation. Cependant, des petites variantes dans la présentation sont possibles, et le système ci-dessous auquel nous ferons référence en est une parmi d’autres.

Nous nous bornerons dans la suite au fragment propositionnel de LL. Il n’y a aucune difficulté à rajouter au cadre que nous présenterons les quantificateurs du premier ordre (cf. par exemple [Girard 91b] et [QuatrTorto 96]), mais nous ne le ferons pas, d’une part parce que le premier ordre ne rajoute rien d’important à notre analyse, et d’autre part parce que nous n’en avons pas parlé jusqu’à présent. Par contre, la gestion du second ordre dans le cas polarisé est plus délicate (parce qu’il faut faire appel à une substitution qui soit elle aussi “polarisée”) : une analyse sémantique a été conduite dans [Quatrini 95], mais (en tous cas dans cette thèse) nous n’en parlerons pas.

Nous aurons besoin dans la section 15.5 de l’unité 1 de LL (le “ V ” de LC et du système LK ci-dessus), ce qui nous conduit à parler des constantes de LL. Si nous avons évité (ou presque) de le faire jusqu’à présent, c’est parce que la version réseaux des constantes additives $\top$ et 0 (en fait $\top$) est loin d’être bien comprise (en tous cas par l’auteur de cette thèse). Les constantes

multiplicatives 1 et $\perp$ sont plus simples à gérer : il faut introduire un lien (logique) l_1 sans prémisses et avec une conclusion typée par la formule 1 , et un lien (logique) $w_\perp$ sans prémisses et avec une conclusion typée par la formule $\perp$. Dans ce dernier cas il faudra étendre la fonction de saut définie dans le chapitre 1 à l'ensemble des liens $w_\perp$. Soyons bien clair : $w_\perp$ se traite *exactement* comme un lien $?w$ (d'où notre notation, qui est un peu abusive). La seule chose à rajouter est que le lien l_1 peut être un saut (c.à.d. qu'il peut être dans le codomaine de la fonction de saut).

Nous allons donc considérer, dans ce chapitre, le calcul des séquents (resp. le système de réseaux) LL_2 (resp. PN_2) propositionnel avec les constantes multiplicatives 1 et $\perp$, auquel nous nous référerons comme à LL_2 (resp. PN_2) propositionnel. Nous donnerons les règles d'introduction de 1 et de $\perp$ dans le cadre du calcul des séquents pour LL_P^f , mais les mêmes règles sont utilisables dans LL_2 . Les e.e.r. sont celles du chapitre 2, auxquelles nous rajoutons l'e.e.r. logique évidente $[1/\perp]$ (il s'agit d'adapter à notre cadre l'e.e.r. de [Girard 87]), la définition de la fonction de saut après cette étape étant la même que celle de la fonction de saut après l'étape $[w]$, *mutatis mutandis*.

Du point de vue sémantique, 1 et $\perp$ sont interprétées par des espaces dont la trame contient un unique élément. Le lecteur n'aura aucun mal à étendre la définition 6.2 d'expérience au "nouveau" système PN_2 propositionnel.

15.2.1. DÉFINITION. (Formules, règles, liens, linéaires (ir)réversibles).

Les règles réversibles (resp. irréversibles) du fragment propositionnel de LL_2 sont les règles introduisant les connecteurs $\wp, \&$ (resp. $\otimes, \oplus$), et la règle introduisant la constante $\perp$ (resp. 1).

Les liens réversibles (resp. irréversibles) du fragment propositionnel de PN_2 sont les liens $\wp, \&, w_\perp$ (resp. $\otimes, \oplus, l_1$).

Toute formule qui est le type d'une arête conclusion d'un lien réversible (resp. irréversible) est dite réversible (resp. irréversible).

15.2.2. DÉFINITION. (Formules linéaires principales). Une (occurrence de) formule dans une preuve π du fragment propositionnel de LL_2 est dite principale lorsque c'est la conclusion d'une règle logique $\wp, \otimes, \&, \oplus$, dérélction, promotion, introduisant $\perp$, introduisant 1 , ou bien la conclusion d'un axiome.

Une (occurrence de) formule dans un réseau R de PN_2 est dite principale lorsque c'est le type d'une arête de R qui est conclusion d'un lien logique $\wp, \otimes, \&, \oplus, ?de, !, w_\perp, l_1$ ou bien d'un axiome.

Lorsque nous souhaiterons être plus précis, nous dirons que la formule A est principale dans une règle (resp. un lien) logique, ou bien que A est principale dans une règle (resp. un lien) axiome.

15.2.3. DÉFINITION. (Formules linéaires polarises). Une formule propositionnelle de LL est dite **positive** (resp. **négative**) et notée P, Q, R (resp. M, N, O) lorsqu'elle est construite de la façon suivante (X étant une formule atomique) :

$$\begin{array}{lcl} P ::= & !X & | \quad P \otimes P \quad | \quad P \oplus P \quad | \quad 1 \quad | \quad !N \\ N ::= & ?X^\perp & | \quad N \wp N \quad | \quad N \& N \quad | \quad \perp \quad | \quad ?P \end{array}$$

avec la convention que seules les formules atomiques peuvent être précédées de deux exponentielles (toute autre formule étant précédée d'au plus une exponentielle).

Une formule de LL est dite polarisée lorsqu'elle est positive, ou bien négative.

La proposition suivante nous dit que si A et B sont équivalentes à $!A$ et $!B$ (resp. $?A$ et $?B$), alors $A \otimes B$ et $A \oplus B$ (resp. $A \wp B$ et $A \& B$) sont équivalentes à $!(A \otimes B)$ et $!(A \oplus B)$ (resp. $?(A \wp B)$ et $?(A \& B)$) : on dira que les connecteurs $\otimes$ et $\oplus$ (resp. $\wp$ et $\&$) préservent le type $!$ (resp. $?$). Nous appellerons “preuve de préservation” du type “!” (resp. “?”) la preuve de $\vdash A^\perp, !A$.

15.2.4. PROPOSITION. Si A est une formule positive (resp. négative) de LL, alors A est prouvablement équivalente (dans LL) à $!A$ (resp. $?A$).

Démonstration : Laissée aux bons soins du lecteur (ou bien cf. [DJS 97]).
□

15.2.5. DÉFINITION. (Preuves et réseaux polariss). Une preuve π de LL_2 propositionnel est polarisée lorsque toutes les formules de π sont des *sous-formules* de formules polarisées.

Un réseau R de PN_2 propositionnel est un réseau polarisé lorsque tous les types de ses arêtes sont des *sous-formules* de formules polarisées.

15.2.6. REMARQUE. Le lecteur nous objectera que toute sous-formule d'une formule polarisée est elle-même polarisée, et il aura *presque* raison : la sous-formule $X^\perp$ de $?X^\perp$ n'est pas polarisée...

La définition ci-dessus dit simplement que nous allons tolérer les réseaux ou les preuves ayant X ou $X^\perp$ parmi leurs conclusions (X étant atomique), comme réseaux ou preuves polarisés.

15.2.7. DÉFINITION. (Preuves et réseaux polarisés et focalisés). Une preuve π polarisée est dite focalisée lorsque toute formule positive qui est active dans une règle irréversible est principale. Nous appellerons LLP l'ensemble de ces preuves.

Un réseau R polarisé est dit focalisé lorsque toute arête ayant comme type une formule positive et qui est prémisses d'un lien irréversible *n'est pas* conclusion d'un lien *coad*. Nous appellerons PNP l'ensemble de ces réseaux.

Remarquons que dans la définition ci-dessus nous nous sommes “amusés” à tourner la phrase définissant un réseau focalisé de façon un peu tordue : en fait, pour une arête ayant comme type une formule positive, ne pas être conclusion d'un lien *coad* dans un réseau c'est forcément être conclusion d'un lien logique ou bien d'un axiome.

Le système qui suit est le calcul des séquents pour LL_P^f : toutes les occurrences de formules sont polarisées, sauf (éventuellement) C . Les séquents Γ, Δ sont des multiensembles de formules polarisées, $?A$ est une formule polarisée, C est sous-formule d'une formule polarisée. Attention ! Les règles irréversibles sont assujetties à la restriction énoncée dans la définition 15.2.7.

LLP

$$\begin{array}{c}
\frac{}{\vdash C^\perp, C} \quad \frac{\vdash \Gamma, P \quad \vdash P^\perp, \Delta}{\vdash \Gamma, \Delta} \\
\\
\frac{\vdash \Gamma, P \quad \vdash \Delta, Q}{\vdash \Gamma, \Delta, P \otimes Q} \quad \frac{\vdash \Gamma, M, N}{\vdash \Gamma, M \wp N} \\
\\
\frac{\vdash \Gamma, M \quad \vdash \Gamma, N}{\vdash \Gamma, M \& N} \quad \frac{\vdash \Gamma, P}{\vdash \Gamma, P \oplus Q} \quad \frac{\vdash \Gamma, Q}{\vdash \Gamma, P \oplus Q} \\
\\
\frac{\vdash ?\Gamma, N}{\vdash ?\Gamma, !N} \quad \frac{\vdash \Gamma, P}{\vdash \Gamma, ?P} \quad \frac{\vdash \Gamma, ?A, ?A}{\vdash \Gamma, ?A} \quad \frac{\vdash \Gamma}{\vdash \Gamma, ?A} \\
\\
\frac{}{\vdash 1} \quad \frac{\vdash \Gamma}{\vdash \Gamma, \perp}
\end{array}$$

15.2.8. REMARQUE. (i) Le lecteur n'aura pas manqué de remarquer que à toute preuve du calcul des séquents LLP est associé un réseau de PNP , et que tout réseau de PNP a une séquentialisation dans LLP (sans que pour autant *toutes* ses séquentialisations soient dans LLP).

(ii) Nous montrerons dans [LaurQuatrTorto 99], que non seulement dans LL_p^f la séquentialisation est beaucoup plus simple que dans LL_2 (comme l'avait remarqué O. Laurent dans [Laurent 99]), mais la normalisation l'est aussi : en utilisant les multiboîtes du chapitre 5 ou bien les réseaux avec poids de [Girard 95a] ou encore les réseaux avec boîtes additives comme nous les avons définis dans le chapitre 1, les restrictions présentes dans LL_p^f permettent de définir très simplement l'e.e.r. $[ccad]$, pour laquelle nous nous sommes donnés tant de mal au chapitre 2.

15.3 LK et LL

Beaucoup de travail a été fait ces dernières années dans la tentative d'extraire un contenu calculatoire des preuves de logique classique. Même si une bonne partie des travaux sur le sujet n'utilisent pas directement la logique linéaire, celle-ci semble jouer un rôle crucial : utilisée comme loupe, d'une part LL suggère une façon raisonnable d'opérer des choix dans la procédure d'élimination des coupures en logique classique, et d'autre part elle permet de classer les solutions proposées jusqu'à présent en mettant en évidence, pour chacune d'entre elles, les choix implicitement opérés, et en agissant de la sorte comme un instrument clarificateur et unificateur (cf. [DJS 97]).

Nous allons résumer brièvement, dans ce paragraphe, l'essentiel des principaux travaux sur le sujet *qui utilisent la logique linéaire*.

La difficulté essentielle dans la compréhension des aspects calculatoires de la logique classique réside dans le non-déterminisme fondamental de la procédure d'élimination des coupures pour les preuves classiques. Il peut y avoir deux façons d'éliminer une même coupure c dans une preuve classique π (c.à.d. que à c sont associées deux étapes élémentaires de calcul) qui sont *foncièrement* différentes : si π_1 et π_2 sont les deux preuves obtenues à partir de π en éliminant la coupure c des deux façons possibles, alors toute sémantique dénotationnelle de la logique classique (qui va identifier π_1 et π_2) identifiera forcément toutes les preuves d'une même formule. On parlera dans la suite de ce phénomène comme d'une "inconsistance algorithmique" du système.

C'est grâce à l'approche de V. Danos, J.-B. Joinet et H. Schellinx (dont l'aboutissement est [DJS 97]) que l'auteur de cette thèse a réellement com-

mencé à comprendre quelque chose aux problèmes évoqués ci-dessus ainsi qu'aux réponses qu'il était possible de fournir en utilisant LL.

Nous allons donc présenter cette approche, fondée sur l'idée de *décoration*.

La première (simplissime) remarque de Danos, Joinet et Schellinx est que si dans une preuve quelconque π de LL_2 , on substitue tout connecteur linéaire par le connecteur classique correspondant, et on efface toutes les exponentielles et les éventuelles répétitions de séquents, on obtient une preuve $sk(\pi)$ de LK (le “squelette” classique de π). On peut alors aussi procéder en sens inverse : une décoration de la preuve α de LK est une preuve $D(\alpha)$ de LL préservant le squelette classique de la preuve, i.e. telle que $sk(D(\alpha)) = \alpha$. On peut alors essayer de définir, sur les preuves de LK , une procédure de normalisation suggérée par les plongements dans LL *qui sont des décorations*, la préservation du squelette nous assurant “qu'on ne s'est pas trop éloigné de la preuve de départ”.

A partir de là, une analyse très générale des décorations possibles d'une preuve donnée est entamée : [Joinet 93], [Schellinx 94], ... Un des aboutissements de ces efforts est que toutes les différentes décorations possibles d'une même preuve classique sont (essentiellement) au nombre de $\dots 2!$ Elles correspondent à deux traductions des formules de LK dans LL : la T -traduction et la Q -traduction.

L'aboutissement de tous ces travaux est [DJS 97] : le point est clairement fait sur les conflits qui apparaissent dans la procédure d'élimination des coupures en logique classique, et des systèmes de calcul viables (construits à partir des traductions T et Q) sont introduits (LK^{tq} , LK^η , LK_{pol}^η).

Une des “surprises” de cette analyse est qu'il n'est pas suffisant (comme l'avait remarqué Girard dans [Girard 91b]) de rajouter une information pour éliminer les conflits dont il était question ci-dessus. Il y a des règles “difficiles” qui ont deux traductions possibles dans LL , irréconciliables entre elles (voir à ce sujet l'amusante annexe B).

Une autre conséquence de cette analyse est que la majorité des systèmes introduits ces dernières années pour extraire un contenu calculatoire des preuves classiques peut être plongée dans LK^{tq} , LK^η ou LK_{pol}^η (notamment FD et le $\lambda\mu$ -calcul de Parigot, cf. [Parigot 91], [Parigot 92] et [JST 95]).

Nous allons définir le système LK_{pol}^η , dans sa version monolatère, ce qui veut dire que notre négation est “très” involutive : on a $A = \neg\neg A$, et les égalités qui s'en suivent par les lois de de Morgan. Ceci est une conséquence de notre approche : puisque la négation est involutive dans LL, nous aurons

un “isomorphisme dénotationnel” (voir aussi le paragraphe 15.5) entre les formules classiques A et $\neg\neg A$.

On pourrait discuter de cette conséquence (a-t-on vraiment $A = \neg\neg A$ en logique classique ?), mais cela nous entraînerait trop loin des objectifs que nous nous sommes fixés, et nous passerons donc la question sous silence. Nous allons ensuite définir la tq -normalisation de [DJS 97], et le P -plongement de Girard de LK_{pol}^η dans LL, pour terminer sur la question que nous résoudrons au paragraphe suivant.

15.3.1. DÉFINITION. Dans le système LK_{pol}^η , à toute formule classique on associe une *couleur* t ou q (couleurs qui sont bien sûr en relation avec les traductions T et Q dont il était question ci-dessus). La couleur d’une formule est déterminée par son connecteur principal (ce qui n’est pas le cas dans les systèmes LK^{tq} et LK^η) : si la règle associée à ce connecteur est réversible (resp. irréversible) dans le sens de la définition 15.1.1, alors la formule est de couleur t (resp. q). Parmi nos formules il y a la constante V et la constante $\neg V$, la première de couleur q et la deuxième de couleur t . La couleur des atomes est choisie arbitrairement, la seule contrainte étant que l’atome X est de couleur t ssi $\neg X$ est de couleur q .

Le calcul des séquents LK_{pol}^η est le calcul LK propositionnel du paragraphe 15.1, mais à présent *toutes* les occurrences de formules sont colorées, Γ, Δ sont des multiensembles de formules colorées, A est une formule colorée. *Et surtout*, les règles irréversibles sont assujetties à la η -**contrainte** : “toute formule de couleur q active dans une règle irréversible est principale”.

15.3.2. REMARQUE. (i) Il découle de la définition ci-dessus que la formule A de LK_{pol}^η est de couleur t ssi la formule $\neg A$ est de couleur q .

(ii) Nous aimerions faire ressentir au lecteur l’élégance de la η -contrainte, mais cela paraît difficile sans en avoir raconté l’histoire et sans en avoir expliqué l’origine (ce que nous ne pouvons faire ici). Signalons tout de même que c’est la première version “abstraite” de la propriété de focalisation que nous ayons rencontrée, et qu’elle permet d’exprimer, en une phrase, ce qu’on était obligé d’écrire auparavant dans les preuves (en utilisant le signe “;” que Girard avait appelé “bénitier” dans [Girard 91b]). La contrainte de la définition 15.2.7 n’est autre qu’une reformulation, dans LL, de la η -contrainte.

Nous décrivons, dans la définition suivante, la procédure d’élimination des coupures définie dans [DJS 97], notamment pour LK_{pol}^η . Nous ferons usage de la notion d’arbre des ancêtres d’une (occurrence de) formule dans une

preuve du calcul des séquents : elle est très naturelle, et nous refusons d'en donner une définition formelle. Toute (occurrence de) formule dans une preuve “provient” d'un axiome d'un affaiblissement ou bien d'une règle logique : en traçant l'“histoire” de la formule on obtient un sous-arbre de la preuve du calcul des séquents, l'arbre dont nous parlons. C'est la notion qui correspond à ce que nous avons appelé (pour les réseaux), dans la définition 7.2.3, le graphe d'une arête.

15.3.3. DÉFINITION. (tq-normalisation) Il y a deux types d'étapes de *tq*-normalisation : les étapes ‘structurelles’ S1 et S2, et les étapes ‘logiques’ L (ou ‘cas-clefs’) :

— On pourra appliquer une étape logique lorsque les prémisses de la coupure c sont toutes les deux principales dans une règle logique. Nous réduirons la coupure c comme un ‘cas-clef’ de Gentzen : on obtient une ou deux coupures résiduelles sur les sous-formules principales de la formule de coupure (suivant le nombre de formules actives dans les règles dont la conclusion est active dans la coupure). Dans le cas de la coupure $\vee_m/\wedge_m$, il est *capital* de remarquer que l'ordre dans lequel on effectue les deux coupures résiduelles est irrélèvent (dans le sens que l'identification sémantique des deux preuves obtenues en inversant cet ordre ne produit pas d'inconsistance algorithmique, voir aussi l'annexe B et pour plus de détails [DJS 97]) grâce à la présence de la η -contrainte. En l'absence de celle-ci, inverser l'ordre de ces deux coupures provoque une inconsistance algorithmique.

— Si aucune étape logique ne peut s'appliquer, alors c'est forcément une étape structurelle qui s'appliquera : celle-ci consiste à ‘transporter’ une des deux sous-preuves dont une conclusion est prémisses de la coupure en ‘remontant’ dans l'arbre des ancêtres de l'autre prémisses de la coupure, en dupliquant la sous-preuve et en contractant les contextes lorsqu'on traverse une règle de contraction (ou le contexte d'une règle binaire additive) ; ce processus termine lorsqu'on atteint des règles axiome, auquel cas les ‘coupures axiome’ qui en résultent sont réduites immédiatement, lorsqu'on atteint des règles affaiblissement, qui sont remplacées par des affaiblissements sur les formules de contexte, ou bien lorsqu'on atteint des règles logiques introduisant le connecteur principal de la formule de coupure.

Il faut (bien sûr !) dire à présent *laquelle* des deux sous-preuves va ‘bouger’. Cela dépendra si oui ou non la prémisses de couleur q de la coupure est principale dans une règle logique. Si c'est le cas, alors c'est la sous-preuve ayant cette occurrence parmi ses conclusions qui va bouger (c'est une étape S_2 , et la coupure est une coupure S_2). Si ce n'est pas le cas, alors ce sera

l'autre (c'est une étape S_1 , et la coupure est une coupure S_1).

Cette procédure diffère essentiellement des procédures 'pas-à-pas' usuelles suivant lesquelles on élimine les coupures en calcul des séquents. Le lecteur n'aura pas manqué de remarquer qu'elle s'inspire de la normalisation des réseaux. Dans [Regnier 92], on définit la 'réduction exponentielle complète', qui ressemble beaucoup à la procédure décrite ci-dessus.

15.3.4. THÉORÈME. *La réécriture définie sur les preuves de LK_{pol}^η par la tq -normalisation est fortement normalisante et confluente.*

Démonstration : Elle fait usage d'un plongement dans les réseaux de PN_2^- , le système dont il a été question dans le chapitre 4. La preuve est dans [DJS 97], et dans une version plus étendue et détaillée dans [JST 95] et [JST 98]. $\square$

Dans la définition qui suit, nous définissons le P -plongement de Girard, pour les formules et les preuves de LK_{pol}^η .

15.3.5. DÉFINITION. (La P -traduction) Le plongement P est défini en utilisant la propriété de préservation du type ? (resp. !) par les connecteurs $\&$ et $\wp$ (resp. $\oplus$ et $\otimes$) énoncée dans la proposition 15.2.4, et en posant pour les atomes $P(X^t) = ?X$ et $P(X^q) = !X$, et pour les constantes $P(V) = 1$ et $P(\neg V) = \perp$.

Le lecteur n'aura alors aucun mal à reconstituer la P -traduction des formules tout seul, mais nous allons tout de même considérer, en guise d'exemple, la traduction d'une formule $A \wedge_m B$ (qui sera de couleur q) :

A	B	$A \wedge_m B$
q	q	$P(A) \otimes P(B)$
q	t	$P(A) \otimes !P(B)$
t	q	$!P(A) \otimes P(B)$
t	t	$!P(A) \otimes !P(B)$

On obtient ainsi une traduction des formules de LK_{pol}^η dans LL , telle que pour toute formule A^q (resp. A^t) $P(A)$ est une formule positive (resp. négative) de LL : on aura en particulier $\vdash P(A)^\perp, !P(A)$ (resp. $\vdash P(A), !P(A)^\perp$).

On définit la traduction du séquent $\vdash \Gamma^t, \Delta^q$ de LK_{pol}^η par le séquent $\vdash P(\Gamma), ?P(\Delta)$ de LL .

La proposition suivante est la proposition 50 de [DJS 97].

15.3.6. PROPOSITION. (Le P -plongement) La P -traduction définie ci-dessus induit un plongement des preuves de LK_{pol}^{η} dans les preuves de LL.

Démonstration : Les règles structurelles et les promotions, utiliserons, à chaque fois que ce sera nécessaire, des coupures avec les preuves de LL de préservation du type ! ou ? des connecteurs. Les règles réversibles sont traduites sans utiliser de règles exponentielles. $\square$

15.3.7. REMARQUE. Le plongement P de LK_{pol}^{η} dans LL n'est pas une décoration (comme nous venons de le voir). Regardons plus précisément pourquoi sur deux exemples :

1) considérons la preuve suivante de LK_{pol}^{η} (où $\Delta^t = \Delta_1^t, \Delta_2^t$ et $\Sigma^q = \Sigma_1^q, \Sigma_2^q$) :

$$\frac{\frac{\pi_1}{\vdots} \quad \frac{\pi_2}{\vdots}}{\frac{\vdash A^t, \Delta_1^t, \Sigma_1^q \quad \vdash B^q, \Delta_2^t, \Sigma_2^q}{\vdash (A^t \wedge_m B^q)^q, \Delta^t, \Sigma^q}}$$

A partir des deux preuves linéaires $P(\pi_1)$ de conclusion $\vdash P(A), P(\Delta_1), ?P(\Sigma_1)$ et $P(\pi_2)$ de conclusion $\vdash ?P(B), P(\Delta_2), ?P(\Sigma_2)$, nous voulons obtenir une preuve de conclusion $\vdash !P(A) \otimes P(B), P(\Delta), ?P(\Sigma)$. Mais pour pouvoir faire une promotion sur $P(A)$, il va falloir (si Δ_1 contient des formules réversibles et pas seulement des formules atomiques) effectuer des déréllections sur le contexte $P(\Delta_1)$. Il faudra ensuite effectuer des coupures avec les preuves de $\vdash !P(A_\delta)^\perp, P(A_\delta)$ (pour les les formules $P(A_\delta)$ de $P(\Delta_1)$ qui ont été déréllectées).

2) considérons la preuve suivante de LK_{pol}^{η} , avec A^t non atomique :

$$\frac{\frac{\pi}{\vdots}}{\vdash A^t, A^t, \Delta^t, \Sigma^q} \quad \vdash A^t, \Delta^t, \Sigma^q$$

On obtiendra une preuve de $\vdash ?P(A), P(\Delta), ?P(\Sigma)$ en effectuant deux déréllections et une contraction à partir du séquent conclusion $\vdash P(A), P(A), P(\Delta), ?P(\Sigma)$ de la preuve $P(\pi)$ de LL. Mais nous voulons une preuve de $\vdash P(A), P(\Delta), ?P(\Sigma)$. Il faudra donc effectuer une coupure avec la preuve de $\vdash P(A), !P(A)^\perp$.

15.4 Le système $LK_{pol}^{\eta, \rho}$

Nous allons introduire, dans ce paragraphe, une restriction sur les preuves de LK_{pol}^{η} , qui fera du plongement P de notre nouveau système ($LK_{pol}^{\eta, \rho}$)

dans LL, une décoration dont l'image sera précisément le système LLP . Nous montrons que $LK_{pol}^{\eta, \rho}$ est un fragment de LK_{pol}^{η} , c.à.d. qu'il est clos par tq -normalisation (théorème 15.4.11), et qu'il est complet par rapport à la prouvabilité classique (théorème 15.4.9).

15.4.1. REMARQUE. L'image de LK_{pol}^{η} à travers le P -plongement est un sous-système de LL (qui n'est pas LLP). Nous avons vu que le plongement P introduit des coupures. Celles-ci sont effectuées entre "l'image de la preuve de départ qu'on a construite inductivement jusqu'au moment d'effectuer la coupure", et une preuve *très* particulière qui est celle de préservation du type "!" ou "?" des formules. Il est donc naturel de se poser la question suivante : que va-t-il se passer si on élimine (dans la preuve de LL) *seulement* les coupures introduites par le plongement ? Va-t-on tomber sur un sous-système LL_S de LL qu'on peut caractériser facilement ? Peut-on découper le sous-système classique S de LK_{pol}^{η} dont les preuves sont les squelettes des preuves de LL_S , comme un fragment (donc stable par tq -normalisation) complet de LK_{pol}^{η} ? (Cette dernière question est très naturelle, si on suit la méthode des décorations).

Nous allons maintenant définir l'opération de renversement d'une formule réversible dans une preuve de LK_{pol}^{η} . On utilisera la notion de η -preuve, définie dans le paragraphe 7.1, dans sa version "classique". La définition qui suit est celle de [DJS 97].

15.4.2. DÉFINITION. Soit $A = B \vee_m C$ (resp. $A = B \wedge_a C$) une formule de LK_{pol}^{η} . Soit π une preuve de conclusion Γ, A de LK_{pol}^{η} . Le renversement de A dans π donne lieu à la preuve π^r obtenue à partir de π de la façon suivante :

1. on substitue, dans π , les axiomes de conclusion $\vdash \neg A, A$ par des η -preuves de même conclusion
2. On efface toutes les règles de π dont la conclusion est la formule A
3. On applique les règles structurelles qui s'appliquaient à A aux formules B et C
4. On rajoute à la preuve ainsi obtenue la règle $\vee_m$ (resp. $\wedge_a$) de conclusion A .

Si $A = \neg V$ et π est une preuve de conclusion Γ, A , alors π^r est obtenue en effaçant toutes les règles introduisant A dans l'arbre des ancêtres de A dans π (c.à.d. en effaçant purement et simplement l'arbre des ancêtres de A dans π), et en rajoutant à la preuve ainsi obtenue la règle introduisant A .

15.4.3. REMARQUE. On peut aussi obtenir la preuve π^r définie ci-dessus, en éliminant la coupure entre π et la η -preuve de $\vdash \neg A, A$ (suivant la *tq*-normalisation). Ce qui montre, bien sûr, que π^r est bien une preuve de LK_{pol}^{η} .

15.4.4. PROPOSITION. Soit π une preuve de LK_{pol}^{η} , et A une formule réversible conclusion de π .

Si π^r est la preuve obtenue en renversant A dans π , alors $[P(\pi)] = [P(\pi^r)]$, où $[P(\alpha)]$ est l'interprétation de la preuve $P(\alpha)$ de LL dans l'espace associé à la conclusion de $P(\alpha)$, pour toute preuve α de LK_{pol}^{η} .

Démonstration : Soit η_A la η -preuve de $\vdash (\neg A)^q, A^t$, et soit $\eta_{P(A)}$ la η -preuve de $\vdash P(A)^\perp, P(A)$.

Notons $cut(\eta_A, \pi)$ (resp. $cut(\eta_{P(A)}, P(\pi))$) la preuve de LK_{pol}^{η} (resp. de LL) obtenue en coupant les preuves η_A et π (resp. $\eta_{P(A)}$ et $P(\pi)$).

Il suffit de remarquer, pour conclure, que la sémantique de la preuve $P(cut(\eta_A, \pi))$ est la même que celle de $cut(\eta_{P(A)}, P(\pi))$. On en déduira que la sémantique de $P(\pi^r)$ (qui est la même que celle de $P(cut(\eta_A, \pi))$ par la remarque ci-dessus) est la sémantique de $P(\pi)$ ($\eta_{P(A)}$ étant interprétée comme la fonction identité de l'espace $\mathcal{P}(A)$). $\square$

Nous introduisons une nouvelle contrainte ρ dans le système LK_{pol}^{η} , qui consiste à exiger que la preuve soit renversée, chaque fois que cela est nécessaire.

15.4.5. DÉFINITION. (La ρ -contrainte)

– *Les contextes réversibles :*

Soit Γ un séquent de LK_{pol}^{η} et soit $A \in \Gamma$. Lorsque le séquent $\vdash \Gamma$ est prémisses d'une règle R dans laquelle A est active, on dit que A est dans un contexte réversible (resp. non réversible) si dans $\Gamma \setminus A$ il y a au moins une (resp. il n'y a aucune) formule réversible.

– *Les règles interdites par ρ :*

Les règles suivantes sont interdites par ρ : les règles structurelles sur les formules réversibles, les règles irréversibles ou les coupures S_1 (plus précisément les coupures dans lesquelles la formule de couleur q n'est principale ni dans une règle logique, ni dans un axiome) dont la prémisses de couleur t est dans un contexte réversible.

– *Les preuves ρ -contraintes :*

Une preuve π de LK_{pol}^{η} est dite ρ -contrainte lorsqu'elle ne contient pas de règles interdites par ρ .

15.4.6. REMARQUE. (i) Le sous-système $LK_{pol}^{\eta, \rho}$ de LK_{pol}^{η} est le système dont les preuves sont les squelettes classiques des preuves de LL obtenues en éliminant les coupures introduites par le P -plongement dans les images par P des preuves de LK_{pol}^{η} .

(ii) Les règles interdites par ρ sont précisément les règles de LK_{pol}^{η} qui requièrent (lors du plongement P) des déréllections sur des formules négatives de LL.

(iii) Nous acceptons le cas de la coupure S_1 dont la prémisse de couleur q est conclusion d'un axiome (quelque soit le contexte de la prémisse de couleur t), ce qui se gère très simplement en traduisant (dans ce cas) l'axiome $\vdash (\neg A)^t, A^q$ par $\vdash P(A)^\perp, P(A)$ au lieu de $\vdash P(A)^\perp, ?P(A)$. On a besoin de ce petit détail dans la preuve de stabilité.

(iv) Remarquons enfin que dans $LK_{pol}^{\eta, \rho}$, toute coupure c de type S_2 est “filiforme” : la formule de couleur q active dans c étant conclusion d'une règle logique, c'est forcément une formule irréversible. Ce qui veut dire que la formule de couleur t active dans c est réversible, et donc par ρ -contrainte ne peut pas avoir été contractée ni affaiblie.

15.4.7. DÉFINITION. (Complexité de réversibilité d'une formule)

Soit A une formule de LK_{pol}^{η} . La complexité de réversibilité de A que l'on note c_A est définie de la façon suivante :

- si $A = \neg V$, $c_A = 1$;
- si A est un atome ou de couleur q , $c_A = 0$;
- si $A = B \vee_m C$ ou $B \wedge_a C$ alors $c_A = c_B + c_C + 1$.

15.4.8. DÉFINITION. (Complexité de réversibilité d'un séquent)

La complexité de réversibilité d'un séquent $\Gamma = \{A_0, \dots, A_n\}$ est égale à 0 si Γ est vide et à la somme des complexités de réversibilité des occurrences de formules qui composent Γ sinon : $c_\Gamma = c_{A_0} + \dots + c_{A_n}$.

Dans la suite (et en particulier dans la preuve du théorème suivant), nous dirons que “ la règle R' suit la règle R dans la preuve π de LK_{pol}^{η} ”, lorsque une des sous-preuves de π prémisses de R' contient la règle R . Nous dirons aussi que “(l'occurrence de) la formule A suit la règle R dans la preuve π de LK_{pol}^{η} ”, lorsque $A \in \Gamma$ et $\vdash \Gamma$ est le séquent conclusion de π ou bien le séquent prémisses d'une règle R' qui suit R dans π . Le lecteur aura donc déjà compris le sens de l'expression “le séquent $\vdash \Gamma$ suit la règle R dans la preuve π de LK_{pol}^{η} ”. D'ailleurs tout cela est complètement trivial, puisqu'une

preuve en calcul des séquents est un arbre définissant un ordre partiel sur les noeuds et les arêtes qui le composent.

15.4.9. THÉORÈME. (Complétude) Soit π une preuve de LK_{pol}^{η} . Il existe, dans LK_{pol}^{η} , une preuve ρ -contrainte π^{ρ} de même conclusion que π .

Démonstration : Par induction sur le nombre de règles interdites par ρ . Soit $n \geq 1$ le nombre de règles interdites par ρ dans π . On prouve qu'il existe une preuve π_1^{ρ} avec au plus $n - 1$ règles interdites. Soit R une règle interdite par ρ et telle que la sous-preuve (resp. les sous-preuves) de π prémisses(s) de R ne contienne(nt) aucune autre règle interdite. Distinguons deux cas :

- si R est une règle structurelle sur une formule réversible A , alors on procède par induction sur la complexité de réversibilité de A : on renverse A dans la sous-preuve de π ayant R comme dernière règle et on applique (éventuellement) l'hypothèse d'induction. Cette transformation a ρ -contraint R , tout en préservant la η -contrainte et sans introduire de nouvelles règles interdites par ρ .

- si R est une coupure S_1 ou une règle irréversible, alors le contexte Γ (resp. l'union des contextes Γ, Γ') de la formule (resp. des formules) de couleur t active(s) dans R contient au moins une formule réversible A . On raisonne par induction sur la complexité de réversibilité de Γ (resp. de $\Gamma \cup \Gamma'$). On a les deux possibilités suivantes :

1. pour toute règle R' qui suit R , R' est irréversible et la formule active dans R' qui suit R est de couleur q ;
2. il existe une règle R' qui suit R et qui n'est pas irréversible, ou bien t.q. la formule active dans R' qui suit R est de couleur t .

Dans le premier cas, on renverse A dans π (qui contient forcément parmi ses conclusions toutes les formules réversibles de Γ (resp. de $\Gamma \cup \Gamma'$)), et on applique (éventuellement) l'hypothèse d'induction.

Dans le deuxième cas, soit R' la "première" règle qui suit R et qui satisfait les hypothèses du cas 2 (c.à.d. la seule règle de π qui satisfait les hypothèses du cas 2 et t.q. toute autre règle de π satisfaisant les hypothèses du cas 2 suit R'). On renverse A dans la sous-preuve de π prémisses de R' et contenant A parmi ses conclusions. Cette transformation a fait diminuer la complexité de réversibilité des séquents prémisses de R , tout en préservant la η -contrainte et sans introduire de nouvelles règles interdites par ρ . $\square$

15.4.10. REMARQUE. Bien sûr, pour une preuve π de LK_{pol}^{η} donnée, la preuve ρ -contrainte obtenue par la méthode ci-dessus décrite, n'est pas unique. Mais les preuves obtenues par cette méthode ne diffèrent entre elles

que par des renversements. Ainsi, une fois plongées dans LL par P , toutes les preuves obtenues ont la même sémantique (par la proposition 15.4.4).

15.4.11. THÉORÈME. (Stabilité) La contrainte ρ dans LK_{pol}^η est stable par élimination des coupures.

Démonstration : Remarquons tout d'abord que la ρ -contrainte (contrairement à la η -contrainte) ne porte que sur les *séquents* prémisses de certaines règles. Ceci implique que si R est une règle de la preuve π de $LK_{pol}^{\eta,\rho}$ qui suit la coupure c , et si π' est obtenue à partir de π en éliminant c , alors R va toujours satisfaire dans π' la contrainte ρ , à moins que R soit devenue dans π' une règle assujétie à la ρ -contrainte alors qu'elle ne l'était pas dans π . La seule possibilité pour que ce soit le cas serait que R soit une coupure de type S_1 dans π' et une coupure d'un autre type dans π . Mais ceci est impossible : le lecteur peut s'en convaincre tout seul, ou bien lire le lemme 2 et la proposition 3 p.230-231 de [JST 95].

On va donc supposer que π est obtenue en effectuant une coupure c entre les preuves λ_1 et λ_2 :

$$\frac{\frac{\lambda_1}{\vdots} \xrightarrow{R_1} \vdash \Gamma, A^t \quad \frac{\lambda_2}{\vdots} \xrightarrow{R_2} \vdash \Delta, (\neg A)^q}{\vdash}^c$$

Soit π' la preuve obtenue en éliminant la coupure c . Nous allons montrer que π' satisfait encore la contrainte ρ . Il faut s'en convaincre pour tous les types possible de coupures c :

- Supposons que c soit une coupure logique, B et C étant les sous-formules principales de A .

Nous allons traiter le cas de la coupure logique $\vee_m / \wedge_m$, en laissant au lecteur le cas $\wedge_a / \vee_a$.

$$\frac{\frac{\frac{1}{\vdots} \xrightarrow{T} \vdash B, C, \dots}{\vdash A^t, \dots}^{R_1} \quad \frac{\frac{2}{\vdots} \xrightarrow{T'} \vdash \neg B, \dots \quad \frac{3}{\vdots} \xrightarrow{T''} \vdash \neg C, \dots}{\vdash (\neg A)^q, \dots}^{R_2}}{\vdash}^c$$

La règle R_1 (resp. R_2) est réversible (resp. irréversible). La preuve π' sera dans ce cas :

$$\begin{array}{c}
\begin{array}{ccc}
1 & 2 & \\
\vdots & \vdots & \\
\hline \vdash B, C, \dots & \vdash \neg B, \dots & \\
\hline \vdash C & & \\
\hline & & \\
\vdash & &
\end{array}
\end{array}$$

$\vdash$

Les seules règles de π' qui ne sont pas des règles de π sont les deux coupures c_1 et c_2 : nous devons vérifier qu'elles sont bien ρ -contraintes. Les seules règles de coupure auxquelles la ρ -contrainte s'applique étant les coupures de type S_1 , nous allons considérer le cas où c_1 (resp. c_2) est une coupure S_1 .

- Supposons que c_1 soit une coupure S_1 . Si $\neg B$ est de couleur q , alors par η -contrainte (dans π) elle est principale dans T' . On a donc deux possibilités :

1. $\neg B$ est principale dans une règle logique : c_1 est alors une coupure S_2 , et ce cas est donc à exclure
2. $\neg B$ est principale dans un axiome : nous sommes alors justement dans le seul cas de coupure S_1 où la ρ -contrainte ne s'applique pas (cf. remarque 15.4.6).

On peut donc supposer que l'occurrence B conclusion de la preuve 1 est de couleur q . La formule $\neg B$ (de couleur t et active dans c_1) est dans le séquent conclusion de T' , et puisque π est ρ -contrainte par hypothèse, dans ce séquent $\neg B$ est dans un contexte non réversible. La coupure c_1 satisfait donc bien la ρ -contrainte.

- Le cas où c_2 est une coupure S_1 se traite exactement de la même façon.

Dans la cas où c est une coupure structurelle, nous allons montrer que toutes les règles dont un des séquents prémisses a changé sont toujours ρ -contraintes. Il ne sera pas nécessaire de vérifier que les coupures éventuellement "créées" par l'étape de normalisation associée à c sont ρ -contraintes, car ce seront forcément des coupures S_2 ou logiques (le lecteur s'en convaincra en relisant la définition 15.3.3), auxquelles la contrainte ρ ne s'applique pas. Remarquons aussi que toute coupure de type S_1 (donc assujétie à la ρ -contrainte) de π' provient forcément d'une coupure de type S_1 de π (toujours par le lemme 2 et la proposition 3 p.230-231 de [JST 95]), ce qui veut dire que toute règle de π' devant satisfaire la ρ -contrainte provient d'une règle de π devant satisfaire la ρ -contrainte.

- Supposons que c soit de type S_2 . La règle R_2 est alors logique introduisant

$(\neg A)^q$. Nous avons vu dans la remarque 15.4.6 que A^t est alors réversible et ne peut avoir été contractée ni affaiblie : la seule feuille de l'arbre des ancêtres de A^t est donc un axiome ou bien une règle réversible T introduisant A^t . La preuve π' est alors la suivante (où dans le cas de l'axiome il faudra substituer la sous-preuve dont la dernière règle est la coupure c' par la preuve λ_2) :

$$\begin{array}{c}
 \begin{array}{cc}
 1 & \lambda_2 \\
 \vdots & \vdots \\
 \hline \vdash A^t, \dots & \vdash (\neg A)^q, \Delta \\
 \hline \vdash \Delta, \dots \\
 \vdots \\
 \hline \vdash \Delta, \dots
 \end{array} \\
 \hline \vdash \Delta, \dots
 \end{array}
 \begin{array}{c}
 \\
 \\
 R_2 \\
 \\
 \\
 \\
 C' \\
 \\
 \\
 \\
 R'_1
 \end{array}$$

Remarquons que R'_1 est la règle R_1 (mais avec des prémisses différentes). Il faut prouver que toutes les règles qui suivent c' sont ρ -contraintes. Cela est très simple, puisque *aucun* des séquents suivant T (ou R_2 dans le cas de l'axiome) dans π' ne peut être assujéti à la ρ -contrainte : dans π , tous les séquents qui suivent T contiennent une formule réversible (A^t) non active dans la règle dont ces séquents sont prémisses.

– Supposons que c soit de type S_1 .

Si $(\neg A)^q$ est conclusion d'un axiome, alors la preuve π' est égale à la sous-preuve λ_1 de π (dont on sait par hypothèse qu'elle est ρ -contrainte). Autrement, la ρ -contrainte s'applique à c dans π : aucune formule de Γ n'est réversible. Il suffira alors de relire la définition de tq -normalisation (15.3.3) pour se rendre compte que cette remarque est bien suffisante pour conclure que π' est ρ -contrainte.

□

15.4.12. PROPOSITION. *Le P -plongement est une décoration de $LK_{pol}^{\rho, \eta}$, dont l'image est contenue dans LLP .*

Démonstration : Il est clair que puisque la contrainte ρ est vérifiée, chaque fois qu'il est nécessaire de faire une promotion en logique linéaire sur la P -traduction d'une formule de couleur t active dans une irréversible ou une coupure S_1 , le contexte est bien modalisé (les formules sont q ou t -atomiques et donc préfixées par un “?”). De plus les règles structurelles ne s'appliquent que sur des formules q ou t -atomiques dans $LK_{pol}^{\eta, \rho}$, et celles-ci sont préfixées par “?” dans la preuve décorée.

Le fait que l'image de $LK_{pol}^{\rho,\eta}$ par P soit précisément LLP est une simple remarque. $\square$

Le plongement P étant une décoration de $LK_{pol}^{\eta,\rho}$, on est “à peu près sûr” que la tq -normalisation pour les preuves de $LK_{pol}^{\eta,\rho}$ peut être “simulée” par la normalisation de LLP . Mais ceci n'est pas tout à fait évident. Nous montrons, dans [LaurQuatrTorto 99], la correspondance exacte entre la tq -normalisation de $LK_{pol}^{\eta,\rho}$, et la normalisation des réseaux de PNP . Cette correspondance a comme conséquence, notamment, la proposition suivante.

15.4.13. PROPOSITION. *Le plongement P de $LK_{pol}^{\eta,\rho}$ dans LLP induit une sémantique dénotationnelle de $LK_{pol}^{\eta,\rho}$: si $\pi \rightarrow_{tq} \pi'$ dans $LK_{pol}^{\eta,\rho}$, alors $[P(\pi)] = [P(\pi')]$, où $[P(\alpha)]$ est l'interprétation de la preuve $P(\alpha)$ de LLP (ou plutôt du réseau de PNP qui lui est associé) dans l'espace associé à la conclusion de $P(\alpha)$, pour toute preuve α de $LK_{pol}^{\eta,\rho}$.*

Le système $LK_{pol}^{\eta,\rho}$ est donc muni d'une procédure d'élimination des coupures (toujours la tq -normalisation) qui a les bonnes propriétés d'un calcul (déduites de celles de LL) et est muni d'une sémantique : l'interprétation d'une preuve π de $LK_{pol}^{\eta,\rho}$ est l'interprétation de $P(\pi)$ dans la sémantique de LL .

15.4.14. THÉORÈME. *Le P -plongement induit une sémantique dénotationnelle pour LK_{pol}^{η} , i.e. si $\pi \rightarrow_{tq} \pi'$ alors $[\pi] = [\pi']$, où, pour toute preuve α de LK_{pol}^{η} , $[\alpha]$ est l'interprétation de $P(\alpha)$ dans l'espace associé au séquent conclusion de $P(\alpha)$.*

Démonstration : On utilise ici la proposition 45 p.53 de [DJS 97] : plaçons-nous dans LK_{pol}^{η} et notons “ $\sim$ ” l'équivalence forte entre preuves (les preuves π et π' sont fortement équivalentes lorsqu'elles diffèrent uniquement par le renversement d'un certain nombre de formules conclusions de π ou actives dans une règle de π), si deux preuves sont fortement équivalentes alors leurs formes normales le sont aussi.

Pour toute preuve α de LK_{pol}^{η} , appelons α^ρ une preuve ρ -contrainte obtenue à partir de α par la méthode utilisée pour prouver le théorème de complétude. On peut vérifier que $\alpha \sim \alpha^\rho$, et on sait (par la proposition 15.4.4) que si $\alpha \sim \alpha'$ alors $[\alpha] = [\alpha']$.

La conclusion est alors une conséquence du résultat et de la remarque ci-dessus, et du fait que P est une sémantique dénotationnelle pour $LK_{pol}^{\eta,\rho}$.

Il suffira de prouver que si π_0 est la forme normale de π , alors $[\pi] = [\pi_0]$. De $\pi \sim \pi^\rho$, on déduit que $\pi_0 \sim \pi_0^\rho$, où π_0^ρ est la forme normale de π^ρ . Puisque P est une sémantique pour $LK_{pol}^{\eta, \rho}$, on a $[\pi^\rho] = [\pi_0^\rho]$ et des remarques initiales on déduit $[\pi_0] = [\pi_0^\rho]$ et $[\pi] = [\pi^\rho]$. D'où $[\pi] = [\pi_0]$. $\square$

15.4.15. REMARQUE. (i) Notre étude permet de retrouver la sémantique des espaces de corrélation qui sont l'ingrédient de base utilisé par J-Y Girard pour construire le système LC ([Girard 91b]). En effet, puisque une preuve de LK_{pol}^η et les preuves ρ -contraintes qui lui sont associées ont la même P -sémantique et que ρ -contraindre signifie transférer les opérations structurelles des formules réversibles sur leurs sous-formules principales, la P -sémantique de LK_{pol}^η revient à considérer que les espaces cohérents associés aux formules réversibles sont munis d'opérations sémantiques de contraction, d'affaiblissement et de contexte de promotion, définies en effectuant ces mêmes opérations sur les sous-formules principales. C'est à dire que l'on interprète les formules réversibles par des *espaces de corrélation négatifs*.
(ii) C'est au moyen de la P -sémantique pour LK_{pol}^η , (le théorème 15.4.14) que nous étudierons des isomorphismes, dans le paragraphe suivant.
(iii) Enfin, la ρ -contrainte permettrait de montrer que la P -sémantique *cohérente multiensembliste* est injective pour le fragment multiplicatif de LC (voir définition 15.5.3), une fois démontrée la conjecture 14.2.10 (voir aussi à ce sujet la remarque 15.5.4).

15.5 Isomorphismes sémantiques et syntaxiques

Dans [Girard 91b], le problème de l'extraction du contenu calculatoire des preuves classiques est posé mathématiquement : il s'agit avant tout de produire une sémantique dénotationnelle pour la logique classique. La question se pose alors de savoir quelles sont les propriétés qu'on va demander à une telle sémantique : nous suffit-il d'obtenir un invariant de l'élimination des coupures ?

La réponse de Girard est qu'il faut chercher une sémantique qui maximise les isomorphismes. Ce qui veut dire que les structures qui interpréteront les formules de la logique classique devront (en tant que structures) satisfaire des isomorphismes tels que la commutativité, l'associativité,... En d'autres termes, on voudrait retrouver les équivalences booléennes comme isomorphismes sémantiques. Si, par exemple, A (resp. B, C) est une formule interprétée par la structure $\mathcal{A}$ (resp. $\mathcal{B}, \mathcal{C}$), on voudra avoir un isomorphisme entre $\mathcal{A} \wedge (\mathcal{B} \wedge \mathcal{C})$ et $(\mathcal{A} \wedge \mathcal{B}) \wedge \mathcal{C}$.

Girard montre qu'une sémantique dénotationnelle pour la logique classique existe, et qu'elle satisfait "beaucoup" d'isomorphismes : c'est la sémantique des espaces de corrélation, une variation sur le thème des espaces cohérents clairement inspirée par un plongement en logique linéaire (le fameux P -plongement des paragraphes précédents). C'est en fait cette sémantique qui le conduit à définir le système de logique classique LC et l'élimination des coupures pour LC .

Du point de vue de LK_{pol}^η , le système LC est assez curieux : c'est un sous-système de LK_{pol}^η , et pourtant il y a une seule conjonction (et pas les deux conjonctions $\wedge_m$ et $\wedge_a$ de LK_{pol}^η) qui est "tantôt multiplicative et tantôt additive" (ce qui reste vrai, dualement, pour la disjonction). En fait, l'image par le plongement P de LC est un sous-système de LL_p^f : celui des formules *fortement polarisées*.

On peut tout d'abord remarquer que les isomorphismes entre espaces de corrélation définis par Girard sont aussi satisfaits par la P -sémantique définie dans le paragraphe précédent. Nous avons ensuite cherché à comprendre le sens, *sur les preuves* de LK_{pol}^η , de ces isomorphismes. Ceci passe par la définition de "isomorphisme syntaxique" qui est non triviale en logique classique. À la lumière de cette définition, les raisons qui ont présidé au choix des styles des connecteurs dans LC apparaîtront très clairement : nous verrons que LC n'est autre qu'une légère amélioration du fragment *multiplicatif* de LK_{pol}^η , ce qui expliquera aussi pourquoi un connecteur mélangeant deux styles différents peut être associatif...

Dans ce paragraphe, les preuves sont moins détaillées qu'à l'ordinaire.

15.5.1. DÉFINITION. (Formules linéaires fortement polarisées). Une formule propositionnelle P (resp. N) de LL est dite **fortement positive** (resp. **fortement négative**) lorsqu'elle est construite de la façon suivante (X étant une formule atomique) :

$$\begin{array}{lcl} P ::= & !X & | \ P \otimes P \quad | \ P \oplus P \quad | \ P \otimes !N \quad | \ !N \otimes P \quad | \ P \oplus !N \quad | \ !N \oplus P \quad | \ 1 \\ N ::= & ?X^\perp & | \ N \wp N \quad | \ N \& N \quad | \ ?P \wp N \quad | \ N \wp ?P \quad | \ ?P \& N \quad | \ N \& ?P \quad | \ \perp \end{array}$$

Une formule de LL est dite fortement polarisée lorsqu'elle est non atomique et *sous-formule* d'une formule fortement positive, ou bien fortement négative.

Une preuve (resp. un réseau) polarisé et focalisé est dit fortement polarisé lorsqu'elle ne contient que (resp. les types de toutes ses arêtes sont) des sous-formules de formules fortement positives ou fortement négatives.

Nous utiliserons pour les formules fortement positives (resp. négatives) les mêmes notations que pour les formules positives (resp. négatives). Lorsque nous ne préciserons pas si nous nous référons à des formules polarisées ou fortement polarisées, cela voudra dire que nos affirmations seront valables dans les deux cas.

15.5.2. REMARQUE. (i) La définition ci-dessus décrète que $?Q$ et $!N$ (avec Q fortement positive et N fortement négative) sont des formules fortement polarisées, sans être ni fortement positives ni fortement négatives.

(ii) L'exemple typique de formule polarisée mais pas fortement polarisée est $?Q \wp ?Q$, avec Q formule fortement positive.

15.5.3. DÉFINITION. (Le système LC et sa P -traduction). LC est le sous-système de LK_{pol}^η obtenu en choisissant une conjonction *additive* entre deux formules de couleurs t , et *multiplicative* autrement. La définition de la disjonction en découle immédiatement par dualité. Nous donnons la P -traduction des formules qui s'en suit :

A	B	$A \wedge_m B$
q	q	$P(A) \otimes P(B)$
q	t	$P(A) \otimes !P(B)$
t	q	$!P(A) \otimes P(B)$
t	t	$P(A) \& P(B)$

La traduction des séquents et des preuves de LC dans LL est celle de LK_{pol}^η dans LL.

15.5.4. REMARQUE. L'image d'une preuve de LC à travers le P -plongement n'est pas une preuve fortement polarisée, et ce pour les mêmes raisons pour lesquelles l'image d'une preuve de LK_{pol}^η à travers le P -plongement n'est pas une preuve polarisée (remarque 15.3.7). Mais, bien sûr, toute preuve du fragment LC^ρ de $LK_{pol}^{\eta,\rho}$ aura comme image par P une preuve fortement polarisée (par définition de LC et par la proposition 15.4.12). C'est la raison pour laquelle la conjecture 14.2.10 permettrait de prouver que la P -sémantique *cohérente multiensembliste* est injective pour le fragment multiplicatif de LC (sans affaiblissements et sans unités).

Venons-en à présent à la définition d'isomorphisme entre deux formules de LK_{pol}^η . Mais... d'abord quelques notations/conventions bien utiles pour la suite.

Pour toute formule A de LK_{pol}^η , on notera ax_A la preuve de LK_{pol}^η ayant comme unique règle l'axiome de conclusion $\vdash \neg A, A$. On notera η_A la η -preuve de même conclusion. Si π (resp. π') est une preuve de LK_{pol}^η ou bien de LL ayant A (resp. $\neg A$ ou $A^\perp$) parmi ses conclusions, nous noterons $cut_A(\pi, \pi')$ la preuve obtenue en effectuant une coupure sur A entre π et π' . Si π est une preuve de LL, on note $[\pi]$ l'interprétation de la preuve π dans une quelconque des trois sémantiques rencontrées jusqu'à présent (cf. chapitre 6).

Commençons par simplifier le problème : plaçons nous dans LL. Comment définir la notion d'isomorphisme entre deux formules de LL ? Rien de plus simple, la voici :

15.5.5. DÉFINITION. (Isomorphisme dans LL). Soient A et B deux formules de LL, et soit ϕ (resp. ψ) une preuve de $\vdash A^\perp, B$ (resp. $\vdash B^\perp, A$). Nous écrirons $A \stackrel{\phi, \psi}{\simeq} B$ lorsque $[\xi_A] = [ax_A]$, $[\xi_B] = [ax_B]$, où $\xi_A = cut_B(\phi, \psi)$ et $\xi_B = cut_A(\psi, \phi)$.

15.5.6. REMARQUE. Rappelons qu'une preuve de LL s'interprète en une fonction (linéaire dans la sémantique cohérente) et qu'une coupure entre deux preuves de LL correspond à la composition de deux fonctions. On peut alors remarquer que, suivant la définition ci-dessus, lorsque $A \stackrel{\phi, \psi}{\simeq} B$ on aura $[cut_B(\phi, \psi)] = [\psi]o[\phi] = id_{\mathcal{A}}$ et $[cut_A(\psi, \phi)] = [\phi]o[\psi] = id_{\mathcal{B}}$, où $\mathcal{A}$ (resp. $\mathcal{B}$) est l'espace interprétant A (resp. B). On aura donc en particulier l'isomorphisme entre les structures interprétant A et B : $\mathcal{A} \simeq \mathcal{B}$.

Pour en venir au cas classique, la définition plus naturelle (eu égard à la définition précédente), en faisant usage de la P -sémantique, est la suivante :

15.5.7. DÉFINITION. Soient A et B deux formules de LK_{pol}^η , et soit ϕ (resp. ψ) une preuve de $\vdash \neg A, B$ (resp. $\vdash \neg B, A$). Nous écrirons $A \stackrel{\phi, \psi}{\longleftrightarrow} B$ lorsque $[P(\xi_A)] = [P(ax_A)]$, $[P(\xi_B)] = [P(ax_B)]$, où $\xi_A = cut_B(\phi, \psi)$ et $\xi_B = cut_A(\psi, \phi)$.

15.5.8. REMARQUE. La définition ci-dessus peut s'exprimer syntaxiquement (sans faire appel au plongement P ni à la sémantique), et c'est ce qui a été fait dans [DJS 96].

15.5.9. REMARQUE. Le lecteur peut vérifier que, suivant cette définition, $A^t \wedge_a B^t \longleftrightarrow A^t \wedge_m B^t$.

Si on prenait la définition ci-dessus comme définition d'isomorphisme, on aurait alors un isomorphisme entre une formule de couleur q (qui se traduit par une formule *positive* de LL) et une formule de couleur t (qui se traduit par une formule *negative* de LL). Il est bien clair que nous n'aurions pas, contrairement au cas de LL (la définition 15.5.5), un isomorphisme entre les structures qui interprètent les deux formules.

C'est en fait la requête qui nous a guidé dans la recherche d'une définition d'isomorphisme syntaxique : quoiqu'on fasse nous voulons que si dans LK_{pol}^η $A \simeq B$, alors dans LL $\mathcal{P}(\mathcal{A}) \simeq \mathcal{P}(\mathcal{B})$.

Regardons de plus près la définition 15.5.7 : qu'est-ce qui s'oppose au fait que si $A \xleftrightarrow{\phi, \psi} B$, alors $\mathcal{P}(\mathcal{A}) \simeq \mathcal{P}(\mathcal{B})$? C'est en fait l'interprétation de la règle de coupure : l'analogue de la remarque 15.5.6 ne s'applique pas car la traduction dans LL d'une règle de coupure *n'est pas* en général (seulement) une coupure de LL.

En logique classique, la règle de coupure semble *ne pas pouvoir* être représentée par une opération mathématique élémentaire. En tous cas, il en est bien ainsi si on s'en tient à l'approche “fonctionnelle” de la logique (une preuve est une fonction) qui est à la base de la correspondance de Curry-Howard, de LL, et d'à peu près tout ce que nous avons dit dans cette thèse. C'est ce que dit la remarque suivante.

15.5.10. REMARQUE. Il n'existe aucune sémantique dénotationnelle de LK_{pol}^η qui puisse interpréter les preuves comme des fonctions et la coupure comme une composition de fonctions.

Démonstration : C'est tout simplement une reformulation de la “paire critique” de Y. Lafont. Considérons la preuve π suivante (nous omettons les contextes par simplicité) :

$$\frac{\frac{\frac{\beta \vdots}{\vdash A^q} \quad \frac{\alpha \vdots}{\vdash (\neg A)^t, B^t} c_1}{\vdash B^t} c_1 \quad \frac{\gamma \vdots}{\vdash (\neg B)^q} c_1}{\vdash} c_1$$

$$\frac{\frac{\beta \vdots \vdots}{\vdash A^q} \quad \frac{\frac{\alpha \vdots \vdots}{\vdash (\neg A)^t, B^t} \quad \frac{\gamma \vdots \vdots}{\vdash (\neg B)^q} c_2}{\vdash (\neg A)^t} c_1}{\vdash} c_2$$

Morale de l'histoire : pour que lorsque $A \simeq B$ dans LK_{pol}^η on ait $\mathcal{P}(A) \simeq \mathcal{P}(B)$, on va demander que la coupure entre les deux preuves ϕ de $\vdash \neg A, B$ et ψ de $\vdash \neg B, A$ dans LK_{pol}^η se traduise par une coupure de LL, c.à.d. que les exponentielles éventuelles apparaissant devant les formules conclusions de $P(\phi)$ et $P(\psi)$ soient *superflues*.

La notion d'exponentielle superflue est assez intuitive : une exponentielle est superflue, lorsque elle est...superflue! (c.à.d. que dans la preuve que nous considérons on aurait pu s'en passer). Nous ne ferons référence que à cette notion très intuitive (que nous définirons d'ici peu) d'exponentielle superflue. Pour une étude approfondie de cette notion, nous renvoyons à [DJS 93].

Nous définissons la notion d'isomorphisme dans LK_{pol}^{η} mais en faisant usage de preuves de $LK_{pol}^{\eta, \rho}$, celles-ci ayant une traduction en réseaux (de PNP) ne faisant pas intervenir de coupures. On pourrait utiliser des preuves de LK_{pol}^{η} , mais cela compliquerait inutilement les choses (la correspondance avec les réseaux étant moins immédiate).

15.5.11. DÉFINITION. Soit $\epsilon \in \{t, q\}$, soit A^ϵ une conclusion de la preuve π de $LK_{pol}^{\eta, \rho}$, et soit R le réseau de PNP associé à π . Appelons a l'arête conclusion de R de type $\mu P(A)$, où $\mu \in \{\emptyset, ?\}$ ($\emptyset$ représente ici la modalité vide).

Nous dirons que A^ϵ est P -principale dans π , lorsque $\mu = \emptyset$, ou bien lorsque $\mu = ?$ et toute arête $b \in G_a^R$ (voir définition 7.2.3 pour la notation) de type $?P(A)$ est conclusion d'un lien *coad* ou bien d'un lien *?de*.

15.5.12. REMARQUE. (et convention) (i) Etre P -principale pour une formule A^q signifie précisément que le ou les lien(s) déréluction dont la conclusion est typée par $?P(A)$ est ou sont “inutile(s)” : il(s) ne ser(ven)t ni pour

effectuer une contraction, ni pour que sa ou leurs conclusion(s) puisse(nt) être (un jour) prémisses d'un lien pax. En d'autres termes, le graphe R^- obtenu à partir de R en éliminant tous ces liens dérélition ainsi que leurs conclusions est un réseau de même conclusion que R , sauf que la conclusion de type $?P(A)$ sera substituée par une conclusion de type $P(A)$. Il est tout aussi clair que les règles dérélition de $P(\pi)$ qui correspondent aux liens ci-dessus peuvent aussi être effacés : on obtiendra de la sorte une preuve π' de LLP, de même conclusion que $P(\pi)$ sauf la conclusion $?P(A)$ qui sera substituée par la conclusion $P(A)$. Nous nous référerons dans la suite à π' comme à $P^-(\pi)$.

(ii) Et bien, nous venons de redécouvrir...le bénitier ! Dire qu'une formule A^q est P -principale dans la preuve π de $LK_{pol}^{\eta,\rho}$, c'est dire (en utilisant le langage de [Girard 91b]) qu'elle est dans le bénitier. Ce qui pourrait induire à se demander si la η -contrainte est véritablement un progrès. Nous n'en doutons pas une seule seconde, mais cela n'empêche pas de représenter les preuves en utilisant le bénitier, ce qui permet de mieux visualiser la formule P -principale. C'est d'ailleurs le choix que nous faisons dans l'annexe B.

(iii) Toute formule de couleur t conclusion d'une preuve π de $LK_{pol}^{\eta,\rho}$ est P -principale dans π .

(iv) Nous adoptons la convention suivante : pour une preuve π de $LK_{pol}^{\eta,\rho}$, nous écrirons $P^-(\pi)$ pour indiquer d'une part qu'il existe une (unique suivant le lemme qui suit) formule conclusion de π de couleur q et P -principale, et d'autre part que $P^-(\pi)$ est la preuve définie en (i) ci-dessus.

(v) Remarquons enfin que si A est P -principale dans π , alors la sémantique de la preuve $P(\pi)$ est toujours égale à celle de la preuve $P^-(\pi)$ à laquelle on a rajouté une règle de dérélition finale dont la formule active est $P(A)$ et la conclusion est $?P(A)$.

Dans le langage de [Girard 91b], A est P -principale dans π lorsque π s'interprète en une *clique centrale* de l'espace de corrélation qui interprète la conclusion de π (voir aussi la proposition 15.5.20).

15.5.13. LEMME. Soit π une preuve de $LK_{pol}^{\eta,\rho}$. Il existe au plus une formule de couleur q qui est P -principale dans π .

Démonstration : Conséquence immédiate de (ii) de la remarque ci-dessus.
□

Nous voici venus à la définition d'isomorphisme.

15.5.14. DÉFINITION. Nous dirons que les formules A et B de LK_{pol}^η sont (ϕ, ψ) -isomorphes, et nous écrirons $A \stackrel{\phi, \psi}{\simeq} B$, lorsque ϕ (resp. ψ) est une preuve de $LK_{pol}^{\eta, \rho}$ de $\vdash \neg A, B$ (resp. $\vdash \neg B, A$), et $P(A) \stackrel{P^-(\phi), P^-(\psi)}{\simeq} P(B)$ dans LL (en fait dans LLP).

15.5.15. REMARQUE. Si $A \stackrel{\phi, \psi}{\simeq} B$, alors les formules de couleur q conclusion de ϕ et ψ sont P -principales (par définition), et ne peuvent donc pas être conclusion de règles structurelles. Remarquons que ce n'est pas le cas lorsque $A \stackrel{\phi, \psi}{\longleftrightarrow} B$ (considérer toujours l'exemple $A^t \wedge_a B^t \longleftrightarrow A^t \wedge_m B^t$).

15.5.16. LEMME. Si $A^\epsilon \simeq B^{\epsilon'}$, alors $\epsilon = \epsilon'$.

Démonstration : C'est évident, par définition d'isomorphisme. $\square$

Le lemme précédent est faux si on substitue l'hypothèse $A^\epsilon \longleftrightarrow B^{\epsilon'}$ à l'hypothèse $A^\epsilon \simeq B^{\epsilon'}$ (comme toujours à cause de $A^t \wedge_a B^t \longleftrightarrow A^t \wedge_m B^t$).

15.5.17. LEMME. Si $A \stackrel{\phi, \psi}{\simeq} B$, alors $A \stackrel{\phi, \psi}{\longleftrightarrow} B$. La réciproque étant fausse.

Démonstration : Il suffit d'écrire (par exemple) la preuve $P(\text{cut}_B(\phi, \psi))$. L'hypothèse de P -principalité va permettre facilement de reconnaître que (modulo des permutations de règles invisibles sémantiquement) ce n'est autre que la preuve $\text{cut}_{P(B)}(P^-(\phi), P^-(\psi))$ avec une déréluction terminale, et les définitions (de $\simeq$ et de $\longleftrightarrow$) permettent alors de conclure.

Concernant la réciproque, on peut reprendre notre exemple favori : on a bien $A^t \wedge_a B^t \longleftrightarrow A^t \wedge_m B^t$. Pourtant on n'a pas $A^t \wedge_a B^t \simeq A^t \wedge_m B^t$, et ce précisément parce que dans la preuve de $\vdash \neg(A^t \wedge_a B^t)^q, (A^t \wedge_m B^t)^q$, la formule $\neg(A^t \wedge_a B^t)^q$ n'est pas P -principale. $\square$

15.5.18. PROPOSITION. Pour $\nu \in \{t, q\}$, on notera $\bar{\nu}$ le seul élément de $\{t, q\}$ différent de ν .

Soit ϕ (resp. ψ) une preuve de $LK_{pol}^{\eta, \rho}$ de $\vdash (\neg A^\epsilon)^{\bar{\epsilon}}, B^{\epsilon'}$ (resp. de $\vdash (\neg B^{\epsilon'})^{\bar{\epsilon}}, A^\epsilon$) sans coupures.

$A^\epsilon \stackrel{\phi, \psi}{\simeq} B^{\epsilon'}$ ssi les trois conditions suivantes sont satisfaites :

- (i) $A^\epsilon \stackrel{\phi, \psi}{\longleftrightarrow} B^{\epsilon'}$
- (ii) $\epsilon = \epsilon'$

- (iii) il existe une règle R (resp. R') introduisant l'unique formule de couleur q conclusion de ϕ (resp. de ψ), telle que toute règle de ϕ (resp. de ψ) suivant R (resp. R') est une règle réversible ou bien une règle structurale (affaiblissement ou contraction) dont la conclusion n'est pas l'unique formule de couleur q conclusion de ϕ (resp. de ψ).

Démonstration : Si $A^\epsilon \stackrel{\phi, \psi}{\simeq} B^{\epsilon'}$, alors la conclusion est une conséquence des lemmes qui précèdent et de la définition de P -principalité.

Réciproquement, la condition (iii) garantit la P -principalité de l'unique conclusion de couleur q de ϕ et de ψ . Comme dans la démonstration du lemme 15.5.17, on pourra alors reconnaître dans la preuve $P(\text{cut}_B(\phi, \psi))$ (modulo des permutations de règles invisibles sémantiquement) la preuve $\text{cut}_{P(B)}(P^-(\phi), P^-(\psi))$ avec une dérélction terminale. La sémantique de cette dernière preuve sera donc celle de $P(\text{cut}_B(\phi, \psi))$, c.à.d. (puisque par hypothèse $A^\epsilon \stackrel{\phi, \psi}{\longleftrightarrow} B^{\epsilon'}$) celle de l'axiome de LL de conclusion $\vdash P(A), P(A)^\perp$ suivi d'une dérélction. Pour conclure que la sémantique de $\text{cut}_{P(B)}(P^-(\phi), P^-(\psi))$ est celle de l'axiome de LL de conclusion $\vdash P(A), P(A)^\perp$ (c.à.d. l'identité de $\mathcal{P}(\mathcal{A})$), on utilise la propriété suivante de la règle de dérélction : si π et π' sont deux preuves de LL de même conclusion dont la dernière règle est une dérélction de conclusion $?A$, si π_1 et π'_1 sont les sous-preuves (de π et π' respectivement) prémisses de cette dérélction, alors lorsque π et π' ont la même sémantique forcément π_1 et π'_1 ont aussi la même sémantique. Nous savons bien, à présent, que c'est loin d'être le cas pour la règle de contraction (cf. chapitre 13)! $\square$

15.5.19. PROPOSITION. Si $A \stackrel{\phi, \psi}{\simeq} B$, alors pour tout connecteur c et pour toute formule D de LK_{pol}^η , $AcD \stackrel{\phi', \psi'}{\simeq} BcD$.

Démonstration : On montre que c'est le cas pour tous les connecteurs. Le point important est que la P -principalité permet d'étendre ϕ et ψ en préservant la η -contrainte : la condition (iii) de la proposition 15.5.18 nous permet d'appliquer la règle irréversible (à laquelle il est impossible d'échapper) correctement. On obtiendra ainsi deux preuves ϕ' et ψ' de $LK_{pol}^{\eta, \rho}$ dont les conclusions de couleurs q seront P -principales. $\square$

15.5.20. PROPOSITION. Si $A \stackrel{\phi, \psi}{\simeq} B$, alors la P -sémantique identifie les deux preuves suivantes (où on a omis les contextes) :

$$\begin{array}{c}
\begin{array}{c} \vdots \\ \vdash A \end{array} \quad \begin{array}{c} \vdots \\ \vdash \neg A \end{array} \\
\hline
\vdash \\
\\
\begin{array}{c} \vdots \\ \vdash A \end{array} \quad \begin{array}{c} \vdots \\ \vdash \neg A, B \end{array} \quad c_1 \quad \begin{array}{c} \vdots \\ \vdash \neg B, A \end{array} \quad \begin{array}{c} \vdots \\ \vdash \neg A \end{array} \quad c_2 \\
\hline
\vdash B \quad \vdash \neg B \quad c_3 \\
\hline
\vdash
\end{array}$$

Démonstration : C'est toujours une conséquence de la P -principalité. Dans le langage de [Girard 91b] une preuve dont la conclusion de couleur q est P -principale est une clique “centrale” dans la sémantique. La propriété que nous énonçons ci-dessus est une conséquence de la centralité des cliques interprétant les preuves ϕ et ψ .

Le lecteur peut aussi se convaincre que la proposition est vraie, tout simplement en écrivant les deux preuves ci-dessus sous forme de réseaux. $\square$

15.5.21. REMARQUE. La proposition ci-dessus est fausse, si on substitue l'hypothèse $A \xrightarrow{\phi, \psi} B$ par $A \xleftarrow{\phi, \psi} B$, et le lecteur peut s'en convaincre en faisant recours à l'exemple habituel.

15.5.22. THÉORÈME. On a $A^t \wedge_m B^t \simeq (A^t \wedge_a B^t) \wedge_m V$.

Démonstration : Il s'agit d'étendre les deux preuves ϕ et ψ telles que $A^t \wedge_a B^t \xleftarrow{\phi, \psi} A^t \wedge_m B^t$, de la seule façon possible pour que les requêtes de P -principalité soient satisfaites. $\square$

15.5.23. REMARQUE. (i) Le théorème ci-dessus n'est autre que l'isomorphisme $!(A \& B) \simeq !A \otimes !B$ de LL. Ce qui n'était pas tout à fait trivial, c'était de l'exprimer dans le langage de la logique classique.

(ii) Ce que dit ce théorème, du point de vue de LK_{pol}^η , c'est que la conjonction multiplicative entre deux formules de couleur t n'est pas une opération “primitive” : elle se décompose en une conjonction additive, et un changement de couleurs.

Le système LC n'est donc rien d'autre que le fragment multiplicatif de LK_{pol}^η , en un peu mieux.

(iii) Puisque la constante V est l'élément neutre du connecteur $\wedge_m$, “dans un contexte multiplicatif, $\wedge_a$ est isomorphe à $\wedge_m$ ”, et c'est bien la raison pour laquelle, malgré un mélange apparent de styles, il existe une unique conjonction *associative* dans LC . Il n'y a en réalité aucun mélange de styles.

Vérifions-le sur un exemple : notons $\wedge$ la conjonction de LC et $\wedge_m$ et $\wedge_a$ (comme d'habitude) les conjonctions, multiplicative et additive, de LK_{pol}^η . On a : $(A^t \wedge B^t) \wedge C^q = (A^t \wedge_a B^t) \wedge_m C^q \simeq (A^t \wedge_a B^t) \wedge_m (V \wedge_m C^q) \simeq ((A^t \wedge_a B^t) \wedge_m V) \wedge_m C^q \simeq (A^t \wedge_m B^t) \wedge_m C^q \simeq A^t \wedge_m (B^t \wedge_m C^q) = A^t \wedge (B^t \wedge C^q)$.
 (iv) L'isomorphisme fondamental $!(A \& B) \simeq !A \otimes !B$ est l'unique "pont" entre le fragment additif et le fragment multiplicatif de LL. C'est ce qui rend les deux fragments non "symétriques", et c'est ce qui explique la raison pour laquelle, en termes d'isomorphismes, il existe une unique solution "optimale" (LC).

Chapitre 16

Discussion : syntaxe et sémantique

Le but de cette discussion est celui de situer les résultats obtenus dans cette thèse dans un cadre moins “strictement technique”.

Notre approche dans les chapitres 6 à 13 ne visait pas à construire une interprétation dans laquelle tout objet du modèle serait interprétation d’une preuve, ce qui serait plutôt une propriété de surjectivité de la sémantique (question qui est connue sous le nom de “full completeness”). Nous voulions montrer que sur les éléments du modèle ayant donné naissance à la logique linéaire, *dont on sait déjà* que ce sont des interprétations de preuves, syntaxe et sémantique coïncident.

Les deux contre-exemples du chapitre 13 sont à cet égard un peu troublants : ils montrent que même dans *MELL*, même pour les cliques *dont on sait déjà* que ce sont des interprétations de preuves, la description intentionnelle des preuves (les réseaux) et leur description extensionnelle (la sémantique cohérente) ne coïncident pas.

Nous avons vu dans le chapitre en question que ceci est à mettre en relation avec la propriété d’uniformité de la sémantique que les réseaux (dans toute leur généralité) ont l’air d’ignorer. Rappelons que cette même uniformité nous a permis de prouver la préservation de l’obsessionnalité (la proposition 9.3.5). Il nous semble qu’il y a ici une voie intéressante à suivre, qui consiste à répondre à la question suivante : “à quoi correspond, en termes de réseaux, la propriété d’uniformité ?” ou encore “existe-t-il une restriction logique ou/et géométrique sur les preuves qui rendrait compte de l’uniformité ?” Plus techniquement, la question que nous nous sommes posé est celle de trouver un sous-ensemble de *MELL* pour lequel la sémantique cohérente

multiensembliste est injective, et de voir si et comment il est possible de caractériser ce sous-ensemble.

Comme le suggère le commentaire qui suit la remarque 13.3.3, nous ne pensons pas que l'uniformité ait quoi que ce soit à voir avec la distinction intuitionniste/classique : nous pensons plutôt qu'elle nous parle de la structure géométrique des preuves.

L'analyse effectuée aux chapitres 6-11, nous a permis de reconduire la question de l'injectivité de la sémantique à celle de l'existence d'une expérience injective pour des réseaux ne contenant que des liens ax , $?de$, $\wp$, $\otimes$, $?w$, $?co$. Dans le chapitre 13 nous montrons que pour certains réseaux non-connexes, une telle expérience n'existe pas. Le lecteur nous objectera que (suivant la définition 1.3.7), tous les graphes de correction sont connexes, et il aura raison. Mais il est bien clair que la fonction de saut définie au chapitre 1 est une façon "ad hoc" de rétablir la connexité, qui permet d'effectuer des "calculs" avec les réseaux, mais qui n'a pas d'équivalent sémantique. La propriété de connexité dont il est question ici est donc celle des graphes de correction sans la fonction de saut (définis dans 1.2.6).

Dans le chapitre 14, nous montrons qu'une expérience injective existe toujours lorsque les réseaux considérés contiennent un graphe de correction (essentiellement) unique qui est connexe (proposition 14.1.4), la connexité étant l'ingrédient *essentiel* dans la démonstration. Ce résultat établi, on peut introduire une "légère dose de non-connexité" (comme nous le faisons dans la proposition 14.2.6) sans perdre la propriété d'existence de l'expérience injective.

Il nous semble qu'il y a là les éléments pour conjecturer une sérieuse relation entre la connexité des réseaux et la propriété d'uniformité de la sémantique cohérente. Ce qui est (pour l'instant) moins clair, c'est comment exprimer cette relation par un énoncé précis.

Plus généralement, les résultats exposés dans cette thèse (et surtout les chapitres 5 et 6-14) s'inscrivent dans la recherche de liens entre syntaxe et sémantique. On voudrait que ces deux approches au calcul "convergent" vers une même représentation des preuves. Il ne s'agit pas tellement (en tous cas au début) d'avoir des résultats "les plus généraux possibles", mais plutôt de comprendre pour quels fragments de LL les résultats que l'on souhaiterait avoir sont vrais.

On peut tenter d'exprimer techniquement les liens étroits entre syntaxe et sémantique que nous cherchons, par le biais d'un certain nombre de propriétés que nous voudrions voir satisfaites. Nous en distinguons trois :

a) surjectivité de la sémantique (full completeness)

b) injectivité de la sémantique

c) confluence de la syntaxe.

La propriété a) est de loin la plus difficile à obtenir. Elle a longtemps mis en défaut les modèles existants (et notamment la sémantique relationnelle et la sémantique cohérente), et a motivé nombre de travaux dans le domaine. Même pour le système *LC* de Girard, la sémantique cohérente multiensemble (ou plutôt sa variante en termes d'espaces de corrélation) n'est pas surjective (cf. [Quatrini 95] pour le contre-exemple).

La recherche d'une solution à ce problème est à la base de la "ludique" : dans ses travaux récents que nous avons déjà cités, Girard modifie la syntaxe et la sémantique de LL, dans le but d'obtenir la propriété a), guidé par les idées qui sont exposées dans [Girard 98a].

Les résultats exposés dans cette thèse peuvent être lus comme une tentative de donner quelques éléments de réponse aux questions b) et c).

Bibliographie

- [AndrPar 91] Andreoli J.-M., Pareschi R., Linear objects : logical processes with built-in inheritance. *New Generation Computing*, 9 (3-4) : 445-473, 1991
- [BarBer 95] Barbanera F., Berardi S., A Symmetric Lambda Calculus for “classical” Program extraction. *Information and Computation*, 125(2) :103-117, 1996.
- [Bechet 94] Bechet D., Second Order Connectives and Proof Transformations in Linear Logic. Presented at the colloquium “Preuves, Réseaux et Types” held in Marseille, France, October 1994, and to appear as a prépublication Université Paris 13, 1999
- [Berry 78] Berry G., Stable models of typed λ -calculi. *Proceedings of ICALP Springer LNCS 62* : 72-89, 1978
- [Danos 90] Danos V., La logique linéaire appliquée à l’étude de divers processus de normalisation (principalement du lambda-calcul). Thèse de doctorat, Université de Paris 7, 1990
- [DJ 99] Danos V., Joinet J.-B., Elementary Recursive Functions in Linear Logic. A paraître dans les actes du colloque “International Workshop on Implicit Computational Complexity” (ICC’99), 1999
- [DJS 95] Danos V., Joinet J.-B., Schellinx H., On the linear decoration of intuitionistic derivations. *Archive for Mathematical Logic*, 33 : 387-417, 1995
- [DJS 93] Danos V., Joinet J.-B., Schellinx H., The structure of exponentials : uncovering the dynamics of linear logic proofs. In Gotlob G., Leitsch A., and Mundici D., editors, *Computational Logic and Proof Theory*, pp. 159-171. Springer-Verlag. *Lecture Notes in Computer Science 713*, Proceedings of the Third Kurt Goedel Colloquium, Brno, Czech Republic, August 1993

- [DJS 97] Danos V., Joinet J.-B., Schellinx H., A new deconstructive logic : linear logic. *Journal of Symbolic Logic*, 62 (3), pp. 755-807, 1997
- [DJS 96] Danos V., Joinet J.-B., Schellinx H., Computational isomorphisms in classical logic. *Electronic Notes in Theoretical Computer Science* 3, et à paraître dans *Theoretical Computer Science*, 1996
- [DanReg 89] Danos V., Regnier L., The structure of multiplicatives. *Archives for Mathematical Logic*, 28, pp. 181-203, 1989
- [DanReg 95] Danos V., Regnier L., Proof-nets and the Hilbert space. In : J.Y. Girard, Y. Lafont and L. Regnier editors, *Advances in Linear Logic*, pp. 307-328. Cambridge University Press, 1995. London Mathematical Society Lecture Note serie 222, Proceedings of the 1993 workshop on Linear Logic, Cornell University, Ithaca.
- [DuqVdW 94] Duquesne E., Van de Wiele J., Modèle cohérent des réseaux de preuve. *Archive for Mathematical Logic* 33 : 133-158, 1994
- [Girard 87] Girard J.Y., Linear logic. *Theoretical Computer Science*, 50 :1-101, 1987.
- [Girard 91a] Girard J.Y., Quantifiers in linear logic II. In : Corsi, G. and Sambin, G. editors, *Nuovi problemi della logica e della filosofia della scienza*, Volume II. Proceedings of the conference with the same name, Viareggio, 8-13 gennaio 1990. CLUEB, Bologna (Italy).
- [Girard 91b] Girard J.Y., A new constructive logic : Classical logic. *Mathematical Structures in Computer Science*, 1(3) : 255-296, 1991
- [Girard 93] Girard J.Y., On the unity of Logic. *Annals of Pure and Applied Logic*, 59 : 201-217, 1992
- [Girard 95a] Girard J.Y., Proof-nets : the parallel syntax for proof-theory. In Ursini and Agliano editors, *Logic and Algebra*, New-York, 1995. Marcel Dekker.
- [Girard 95b] Girard J.Y., Light Linear Logic. *Information and Computation*, vol. 143, n. 2, p.175-204, 1998
- [Girard 98a] Girard J.Y., On the meaning of logical rules 1 : syntax vs. semantics. In U. Berger and H. Schwichtenberg, editors, *Computational Logic*, Heidelberg, 1999. Springer-Verlag. NATO series F 165.

- [Girard 98b] Girard J.Y., On the meaning of logical rules 2 : multiplicatives and additives. Prépublication, Institut de Mathématiques de Luminy, Marseille, 1998
- [Joinet 93] Joinet J.-B., Etude de la normalisation du calcul des séquents classique à travers la logique linéaire. Thèse de doctorat, Université Paris 7, 1993
- [JST 95] Joinet J.-B., Schellinx H., Tortora de Falco L., Strong Normalization for “all-style” LK^{tq} . Extended abstract, Proceedings of the conference “Theorem Proving with Analytic Tableaux and Related Methods”, 5th international Workshop, Tableaux 96, Terrasini, Palermo Italy, May 1996, LNAI 1071 pp.226-243, Springer-Verlag 1996.
- [JST 98] Joinet J.-B., Schellinx H., Tortora de Falco L., SN and CR for free-style LK-tq : linear decorations and simulation of normalization. Preprint nr. 1067, May 1998 Department of Mathematics, Universiteit Utrecht
- [Joly 00] Joly Thierry, Codages, séparabilité et représentation de fonctions en λ -calcul simplement typé et dans d’autres systèmes de types. Thèse de doctorat, Université Paris 7, 2000
- [KriPar 90] Krivine J.-L., Parigot M., Programming with proofs. Journal of Information Processing and Cybernetics EIK (formerly Elektronische Informationsverarbeitung und Kybernetik) 26, 1990
- [Laurent 99] Laurent Olivier, Polarized Proof-Nets : Proof-Nets for LC (Extended Abstract). In Jean-Yves Girard, editor, Typed Lambda Calculi and Applications ’99, volume 1581 of LNCS, pages 213-227. Springer-Verlag.
- [LaurQuatrTorto 99] Laurent O., Quatrini M., Tortora de Falco L., Logique linéaire polarisée et logique classique. En préparation, 1999
- [Parigot 91] Parigot M., Free Deduction : an analysis of computation in classical logic. In : Voronkov, A, editor, Russian Conference on Logic Programming, pp. 361-380. Springer Verlag. LNAI 592. 1991.
- [Parigot 92] Parigot M., $\lambda\mu$ -calculus : an algorithmic interpretation of classical natural deduction. In Voronkov A. editor, Logic Programming and Automatic Reasoning, pp. 190-201. Springer-Verlag. LNAI 624, 1992

- [Quatrini 95] Quatrini Myriam, Sémantique cohérente des exponentielles : de la logique linéaire à la logique classique. Thèse de doctorat, Université Aix-Marseille II, 1995
- [QuatrTorto 96] Quatrini Myriam, Tortora de Falco L., Polarisation des preuves classiques et renversement. Compte-rendu de l'Académie des Sciences, t. 323, Série I, pp. 113-116, 1996
- [Regnier 92] Regnier L., λ -calcul et réseaux. Thèse de doctorat, Université Paris 7, 1992
- [Schellinx 94] Schellinx Harold, The Noble Art of Linear Decorating. ILLC Dissertation Series, 1994-1. Institute of Logic, Language, and Computation, University of Amsterdam, 1994
- [Statman 83] Statman R., λ -definable functionals and $\beta\eta$ conversion. Arch. Math. Logik Grundlagenforsch. 23, 21-26, 1983
- [Torto 97] Tortora de Falco L., Denotational semantics for polarized (but non-constrained) LK by means of the additives. Lecture Notes in Computer Science, 1289 pp. 290-304 special issue for the proceedings of the conference "5th Kurt Goedel Colloquium KGC'97, computational logic and proof theory", Wien (Austria), 25-29 août 1997
- [Torto 00a] Tortora de Falco L., Additives of Linear Logic and Normalization -Part I : a (restricted) Church-Rosser Property. A paraître dans Theoretical Computer Science, Special issue for the Proceedings of the conference "Linear Logic 96" held in Tokyo april 1996
- [Torto 00b] Tortora de Falco L., Additives of Linear Logic and Normalization -Part II : the additive standardization Theorem, soumis à la revue Theoretical Computer Science

Annexe A

Le théorème de standardisation additive

La question qui nous occupera dans cette annexe est celle de la normalisation forte des réseaux de PN_2 . La méthode des *candidats de réductibilité* introduite par Girard pour démontrer la normalisation forte du système F est appliquée à la logique linéaire dans [Girard 87]. Mais, comme nous l'avons déjà remarqué dans la démonstration du théorème 4.0.5, la syntaxe linéaire rend indispensable la preuve d'un résultat supplémentaire, connu sous le nom de "théorème de standardisation". Dans [Danos 90], ce résultat est démontré pour le sous-système PN_2^- de PN_2 , ne contenant pas les connecteurs additifs. Il s'agissait de prouver ce même résultat pour les additifs.

Les connaissances acquises dans la première partie de la thèse (et notamment le théorème 3.2.2) nous permettent de prouver le théorème de standardisation additive pour PN_2 (théorème A.2.2).

Il faut dire, cependant, que ce théorème ne permet pas de conclure que le système PN_2 est fortement normalisant (voir remarque A.5.2). Si les résultats à notre disposition semblent être suffisants pour conclure, nous n'avons toujours pas à ce jour une preuve complète de ce théorème. Il n'est peut-être pas inutile de signaler qu'une telle démonstration n'existe nulle part.

What we refer to as "the independence lemma" in the sequel, is theorem 3.2.2. We start by introducing a convention.

Convention : Suppose that $T \xrightarrow{\sigma} T'$ and that $\mathbb{B}$ is an additive box with exponential depth 0 in T . If $\mathbb{B}'$ is a residue of $\mathbb{B}$ in T' with additive depth $p > 0$, then there exist p additive boxes $B_1, \dots, B_p$ of T' different from $\mathbb{B}'$

and s.t. $\mathbb{B}'$ is a subproof-net of $B_i \forall i \in \{1, \dots, p\}$. In the sequel we will speak of the **main subproof-nets** of T' (that will also be denoted by $\beta_1, \dots, \beta_p$) in the canonical representation of T' with respect to $\mathbb{B}'$, meaning the components of the additive boxes $B_1, \dots, B_p$ which do not contain $\mathbb{B}'$.

The independence lemma states that all the residues of $\mathbb{B}$ different from $\mathbb{B}'$ are subproof-nets of a main subproof-net of T' in the canonical representation of T' w.r.t. $\mathbb{B}'$.

We first prove the following useful lemma

A.0.24. LEMMA. Let T, T' be two proof-nets such that $T \xrightarrow{\sigma} T'$, let $\mathbb{B}$ be an additive box with exponential depth zero in T and $\mathbb{B}'$ a residue of $\mathbb{B}$ in T' . Given a proof-net R appearing in σ , we will refer to the canonical representation of R with respect to the (unique) lift of $\mathbb{B}'$ in R .

Let T_1 and T_2 be two proof-nets in the following decomposition of $\sigma : T \xrightarrow{\sigma_1} T_1 \xrightarrow{\sigma_2} T_2 \xrightarrow{\sigma_3} T'$. If β is a main subproof-net of T_1 , then all its residues in T_2 are subproof-nets of main subproof-nets of T_2 .

proof : By induction on σ_2 .

If $T_1 = T_2$, then the lemma holds.

Otherwise, $T_1 \xrightarrow{\sigma'_2} T'_2 \xrightarrow{[x]} T_2$. If β has no residue in T_2 then the lemma vacuously holds. Otherwise, let β be a main subproof-net of T_1 , β'' a residue of β in T_2 and β' its lift in T'_2 .

By induction hypothesis, β' is a subproof-net of a main subproof-net γ of T'_2 .

If γ has no residue in T_2 , then $\gamma \xrightarrow{[x]} \delta$ (γ cannot disappear in a $[\&/\oplus]$ e.r.s., because in this case β' would disappear too) where δ is a main subproof-net of T_2 , and β'' is a subproof-net of δ .

Otherwise, β'' is a subproof-net of a residue of γ in T_2 . We are now going to show that for every possible $[x]$, any residue of γ in T_2 is a subproof-net of a main subproof-net of T_2 .

Let $\mathbb{B}'_2$ (resp. $\mathbb{B}_2$) be the lift of $\mathbb{B}'$ in T'_2 (resp. T_2). Let B' be the additive box of T'_2 s.t. γ is a component of B' . By definition of main subproof-net, $\mathbb{B}'_2$ is a subproof-net of the component of B' different from γ .

Observe that B' has one or two residues in T_2 (otherwise γ or $\mathbb{B}'_2$ would have no residue in T_2). Let's now consider these two possible cases.

If B' has one residue B in T_2 , then $\mathbb{B}_2$ (the unique residue of $\mathbb{B}'_2$ in T_2) and the unique residue of γ in T_2 will (still) be subproof-nets of the two different

components of B . To conclude, remember that the component of B which does not contain $\mathbb{B}_2$, i.e. the one containing γ , is a main subproof-net of T_2 . If B' has two residues B_1 and B_2 in T_2 , then $[x] = [ccad]$ and B' is a subproof-net of the subproof-net of T'_2 duplicated by $[x]$. This means that there are two residues of γ : the component γ_1 of B_1 and the component γ_2 of B_2 . Let B'' be the box of T'_2 s.t. $[x]$ is B'' -attractive, and let's call B_3 the unique residue of B'' in T_2 . $\mathbb{B}_2$ is either a subproof-net of B_1 or a subproof-net of B_2 . Suppose for example that $\mathbb{B}_2$ is a subproof-net of B_1 : then γ_1 is the component of B_1 which does not contain $\mathbb{B}_2$ (i.e. γ_1 is a main subproof-net of T_2), and γ_2 (is a subproof-net of B_2 which) is a subproof-net of the component of B_3 which does not contain $\mathbb{B}_2$ (and this component of B_3 is a main subproof-net of T_2). $\square$

A.1 Normalization without interaction

Let Q be a proof-net having the edge a labeled by A among its conclusions, let S be a proof-net having the edge $a^\perp$ labeled by $A^\perp$ among its conclusions, and let T be the proof-net obtained from Q and S by connecting the edges a and $a^\perp$ with a cut link c . In the whole section, Q, S, T will denote proof-nets of the previous form, and c will denote the cut link between Q and S .

Suppose that $T \xrightarrow{\sigma} T'$ and that for every $[x] \in \sigma$, $[x]$ is not a $[ccad]$ e.r.s. and x is not a residue of c . Then the reader will easily convince herself/himself that there exists a subproof-net Q' of T' and a subproof-net S' of T' , s.t. $Q \rightarrow Q', S \rightarrow S'$, and if we connect Q' and S' by means of a cut link with premises the residue of a in Q' and the residue of $a^\perp$ in S' , then we obtain T' . Intuitively, this means that until we don't reduce a residue of c , Q and S will never interact.

The main point of this section (proposition A.1.7) is to show that even though the previous statement is wrong if we allow $[ccad]$ e.r.s., its intuitive meaning is still true. This result, as well as the five lemmas closing the section, will be crucial in section A.4. The reader who "believes" that these results hold can skip this section.

A.1.1. REMARK. Up to the renaming of bound variables, the label A (resp. $A^\perp$) of the conclusion a (resp. $a^\perp$) of Q (resp. S) does not contain any occurrence of a free variable X s.t. X is bound by a forall link somewhere in T . This implies that if $T \xrightarrow{\sigma} T'$, then the two premises of every residue of c in T' are still labeled by A and $A^\perp$.

The following lemma says (in particular) that under certain circumstances the conditions (1) and (2) of remark 2.1.3 are stable w.r.t. cut elimination.

A.1.2. LEMMA. *Let R, R' be two proof-nets such that $R \xrightarrow{[x]} R'$. Let n be a main link of R and Φ a direct path of R with starting link n .*

(i) *Suppose that Φ has the cut link c_1 as terminal link, and suppose that c_1 has a residue in R' . For every residue n' of n in R' , there exists a residue c' of c in R' , and a direct path Φ' having n' as starting link and c' as terminal link.*

(ii) *Suppose that Φ has the conclusion a of R as terminal edge. For every residue n' of n in R' , there exists a direct path Φ' having n' as starting link and the residue a' of a in R' as terminal edge.*

proof : Let n' be a residue of n in R' (if there is no such link, then the lemma vacuously holds). We proceed by inspection of all the possible cases for $[x]$.

If x is a logical cut link : then n' is the unique residue of n in T' .

- in case (i) : $x \neq c_1$ and c_1 has a unique residue c' in R' . The direct path Φ' of R' is then the same as the path Φ of R , except when x is a $\&/\oplus$ or a $!/?de$ cut link and n belongs to the box opened by $[x]$ while c_1 doesn't : in this case Φ' crosses a coad or a pax link less than Φ .

- in case (ii) : the direct path Φ' of R' is the same as the path Φ of R , except when x is a $\&/\oplus$ or a $!/?de$ cut link and n belongs to the box opened by $[x]$: in this case Φ' crosses a coad or a pax link less than Φ .

If $[x]$ is an axiom e.r.s. : then n' is the unique residue of n in T' .

- in case (i) : $x \neq c_1$ and c_1 has a unique residue c' in R' . The direct path Φ' of R' is then the same as the path Φ of R .

- in case (ii) : the direct path Φ' of R' is the same as the path Φ of R .

If $[x]$ is a contraction or a $[ccad]$ e.r.s. :

- in case (i), we still have to distinguish four cases :

(i.1) c_1 is duplicated by $[x]$ (which means that $[x] = [c_1]$ or that c_1 belongs to the subproof-net of R duplicated by $[x]$) and n is also duplicated by $[x]$ (which means that n belongs to the subproof-net of R duplicated by $[x]$) : then whatever of the two residues of n one chose for n' , the direct path Φ' of R' is the same as the path Φ of R and c' is the appropriate residue of c_1 in R' .

(i.2) c_1 is duplicated by $[x]$ and n isn't : then $[x] = [c_1]$. If $[x]$ is a contraction e.r.s., then c' is the appropriate residue of c_1 in R' and Φ' crosses one contraction link less than Φ . If $[x]$ is a $[ccad]$ e.r.s., then c' is the appropriate residue of c_1 in R' and Φ' crosses one coad link less than Φ .

(i.3) n is duplicated by $[x]$ and c_1 isn't : then $x \neq c_1$ and c_1 has a unique residue c' in R' . If $[x]$ is a contraction e.r.s., then Φ' crosses one contraction link more than Φ . If $[x]$ is a $[ccad]$ e.r.s., then Φ' crosses one coad link more than Φ .

(i.4) Neither n nor c_1 is duplicated by $[x]$: then n and c_1 both have a unique residue in R' and Φ' is the same as Φ .

- in case (ii) we still have to distinguish two cases :

(ii.1) n is duplicated by $[x]$ (which means that n belongs to the subproof-net of R duplicated by $[x]$) : then Φ' crosses one contraction (resp. one coad) link more than Φ if $[x]$ is a contraction (resp. a $[ccad]$) e.r.s..

(ii.2) n is not duplicated by $[x]$: then Φ' is the same as Φ .

If x is a weakening cut link : then n cannot belong to the exponential box erased by $[x]$ (we are supposing that n has a residue n' in R'). Let m be the weakening link of R whose conclusion is one of the premises of x . If $n \neq m$, then in both cases ((i) and (ii)) Φ' is the same as Φ . If $m = n$, then we should be in case (i) and we should have $c_1 = x$, but this is impossible because c_1 would have no residue in R' .

If x is a cc cut link : then n has a unique residue in R' and (in case (i)) c_1 has also a unique residue c' in R' . Let $B_1^!$ be the exponential box of R swallowed by the exponential box $B_2^!$ of R when performing $[x]$. If n does not belong to $B_1^!$ nor $B_2^!$, then Φ' is the same as Φ . Otherwise :

- in case (i) : if n belongs to $B_1^!$ and $c_1 = x$ or if n and c_1 belong to $B_1^!$, then Φ' is the same as Φ . If n belongs to $B_1^!$ and $c_1 \neq x$ doesn't, then Φ' crosses one pax link more than Φ . If n and c_1 belong to $B_2^!$ or if n belongs to $B_2^!$ and $c_1 \neq x$ doesn't, then Φ' is the same as Φ . If n belongs to $B_2^!$ and $c_1 = x$, then Φ' crosses one pax link less than Φ .

- in case (ii) : if n belongs to $B_1^!$ then Φ' crosses one pax link more than Φ , and if n does not belong to $B_1^!$ then Φ' is the same as Φ . $\square$

A.1.3. LEMMA. Let R, R' be two proof-nets s.t. $R \xrightarrow{[x]} R'$. Let n_1, n_2 be two main links of R s.t. $n_1 \neq n_2$, and let c_3 be a cut link of R s.t. if $[x] \neq [ccad]$, then $x \neq c_3$. Let n'_1 (resp. n'_2) be a residue of n_1 (resp. n_2) in R' and let c' be a residue of c_3 in R' .

Suppose that there exists a couple (Φ_1, Φ_2) of direct paths of R s.t. Φ_i has n_i as starting link and c_3 as terminal link. Suppose also that there exists a couple (Φ'_1, Φ'_2) of direct paths of R' s.t. Φ'_i has n'_i as starting link and c' as terminal link. If $\Phi_1 \cap \Phi_2 = \{c_3\}$, then $\Phi'_1 \cap \Phi'_2 = \{c'\}$.

proof : Clearly $c' \in \Phi'_1 \cap \Phi'_2$. Φ'_1 and Φ'_2 are nothing but the direct paths associated with Φ_1 and Φ_2 given by (i) of the previous lemma. We know

that $\Phi_1 \cap \Phi_2 = \{c_3\}$ and we have to show that the e.r.s. $[x]$ does not change this situation. We just explain why the lemma holds : from the proof of (i) of the previous lemma, one can see that for $i \in \{1, 2\}$ either $\Phi'_i = \Phi_i$, or Φ'_i crosses one contraction, coad or pax link more or less than Φ_i . It is then easy to see that, in any case, none of the links created by $[x]$ can be crossed by both Φ'_1 and Φ'_2 (the hypothesis $n_1 \neq n_2$ is of course crucial).

A detailed proof can be given by inspecting (again!) all the possible cases for $[x]$. $\square$

Let $T \xrightarrow{\sigma} T'$. Every main link (see lemma 2.1.4) of T' is a residue of a link of Q or of a link of S . The following definition introduces a special terminology for these links, and gives also a similar “status” to the links of T' which are not main links.

A.1.4. DEFINITION. Suppose that $T \xrightarrow{\sigma} T'$. Let n' be a link of T' , and let m' be a main link of T' s.t. there exists a direct path with starting link m' and terminal link n' . The existence of m' is given by lemma 2.1.4 (if n' is a main link, then $m' = n'$). Let m be the lift of m' in T .

We say that m' is a **main Q -link** (resp. a **main S -link**) iff m is a link of the subproof-net Q (resp. S) of T . We say that n' is a **Q -link** (resp. an **S -link**) iff m' is a main Q -link (resp. a main S -link). The conclusion of an S -link (resp. a Q -link) is called an **S -edge** (resp. a **Q -edge**). An additive box B is called a **Q -box** (resp. an **S -box**) iff its front door is a main Q -link (resp. a main S -link).

We say that n' is a **QS -link** (or an **SQ -link**) if n' is both an S -link and a Q -link. We use the same terminology for edges.

A.1.5. REMARK. Let T' be a reduct of T .

- (i) Any main link of T' is a main Q -link or a main S -link and cannot be both.
- (ii) Let Φ and Ψ be two direct paths of T' s.t. the starting link of Φ (resp. Ψ) is a Q -link (resp. an S -link). If m is a link of T' s.t. $m \in \Phi \cap \Psi$, then m is a QS -link.
- (iii) The unique link of T which is a QS -link is the cut link c .

The following lemma is obviously wrong if we substitute to B^1 an additive box (suppose for example that $[c]$ is a $[ccad]$ e.r.s. and that $T \xrightarrow{[c]} T'$).

A.1.6. LEMMA. Let $T \xrightarrow{\sigma} T'$ and let $B^\dagger$ be an exponential box of T' . If for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a QS -link, then one (and only one) of the two following conditions holds :

- (1) every link of $B^\dagger$ is a Q -link which is not an S -link
- (2) every link of $B^\dagger$ is an S -link which is not a Q -link.

proof : By induction on σ . If $T' = T$ the result is obvious. Otherwise $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$. By induction hypothesis, the result holds for the exponential boxes of T_1 . Suppose now that $B^\dagger$ does not satisfy (1) nor (2) of our lemma, let n' be the of course door of $B^\dagger$ and let n_1 be its lift in T_1 .

If n' is a main Q -link, then n_1 is a main Q -link. By induction hypothesis the exponential box $B_1^\dagger$ of T_1 associated with n_1 contains then only Q -links that are not S -links. The only possibility to have a different situation for $B^\dagger$ in T' is that some S -links enter the box $B_1^\dagger$ when performing $[x]$. Then $[x]$ must be a $[cc]$ e.r.s. s.t. an exponential box $B_2^\dagger$ of T_1 (containing some S -links) is swallowed by $B_1^\dagger$. But if $B_2^\dagger$ contains an S -link, then (by induction hypothesis) every link of $B_2^\dagger$ is an S -link. On the other hand the pax doors of $B_1^\dagger$ are Q -links, so that x is a QS -link of T_1 , thus contradicting the hypothesis on σ .

If n' is a main S -link the proof is obviously the same. $\square$

We have seen that the unique link of T which is a QS -link is the cut link c . Roughly speaking, the following proposition states that until one doesn't reduce any residue of c , this is still the case for any reduct of T . We give a detailed proof of this result.

A.1.7. PROPOSITION. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$, x is not a residue of c .

If the link n of T' is a QS -link, then it is a residue of c in T' . (In particular, there are no QS -edges in T').

proof : By induction on σ . If $T = T'$, then the result is obvious. Otherwise $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$. There is no link of T' which is a main Q -link and a main S -link, and this means that there exist a main Q -link n'_1 of T' (resp. a main S -link n'_2 of T') and a direct path Ψ'_1 (resp. Ψ'_2) with starting link n'_1 (resp. n'_2) and terminal link n . $n'_1 \neq n'_2$. Let then n_1 (resp. n_2) be the lift of n'_1 (resp. n'_2) in T_1 : by induction hypothesis n_1 is a main Q -link which is not an S -link, and n_2 is a main S -link which is not a Q -link. In particular $n_1 \neq n_2$.

Part I : We first show that (2) of remark 2.1.3 cannot hold for n .

If (2) holds for n , then (2) holds for n'_1 and n'_2 . Let a' and Φ' be the conclusion and the direct path of T' associated with n given by remark 2.1.3. Let a_1 be the lift of a' in T_1 .

If (2) of remark 2.1.3 holds for n_1 and for n_2 , then from lemma A.1.2, we deduce the existence of a direct path Φ_1 (resp. Φ_2) with starting link n_1 (resp. n_2) and terminal edge a_1 . Then a_1 is a QS -edge, thus contradicting the induction hypothesis.

We can then suppose (for example) that (1) of remark 2.1.3 holds for n_1 . Let c_1 and Ψ_1 be the cut link and the direct path of T_1 associated with n_1 given by remark 2.1.3. By lemma A.1.2 c_1 has no residue in T' . Because n'_1 is a residue of n_1 in T' , c_1 cannot belong to a subproof-net of T_1 erased by $[x]$. Then (from remark 2.2.10) $x = c_1$ is a logical cut link (case (i)) or $x = c_1$ is a weakening cut link (case (ii)), or $[x] = [c_1]$ is an axiom e.r.s. (case (iii)).

In case (i), clearly (1) of remark 2.1.3 would still hold for n'_1 , and this cannot be the case.

Case (ii) is represented in figure 14. n_1 is necessarily the weakening link whose conclusion is a premise of x . n_1 is a main Q -link and c_1 is then a Q -link. Because $[c_1] \in \sigma$ and $[c_1] \neq [ccad]$, c_1 is not a residue of c in T_1 , and then (by induction hypothesis) c_1 is not an S -link. Let $B_1^!$ be the exponential box erased by $[c_1]$. The induction hypothesis applied to all the proof-nets appearing in σ_1 allows to deduce that for every $[y] \in \sigma_1$ s.t. $[y] \neq [ccad]$, y is not a QS -link : otherwise y would be a residue of c and we would have $[y] \in \sigma$, thus contradicting the hypothesis on the reduction sequence σ . We can then apply lemma A.1.6 to $B_1^!$: the conclusion of the main door of $B_1^!$ is a Q -edge which is not an S -edge (because c_1 is not an S -link) and then every link of $B_1^!$ is a Q -link which is not an S -link.

The S -link n_2 cannot be a link of $B_1^!$, and then (2) of remark 2.1.3 holds for n_2 : if (1) would hold for n_2 , then the cut link c_2 of T_1 given by remark 2.1.3 would necessarily have a residue in T' and (from lemma A.1.2) (1) would hold for n'_2 . Lemma A.1.2 gives the existence of a direct path Ψ_2 (represented in figure 14) of T_1 with starting link n_2 and terminal edge a_1 . The fact that (2) of remark 2.1.3 holds for n'_1 implies the existence of the path Ψ (represented in figure 14) with starting link the pax door m_1 of $B_1^!$ and terminal edge a_1 . m_1 is a Q -link while n_2 is a (main) S -link, so that figure 14 clearly shows that a_1 is a QS -edge, which contradicts the induction hypothesis.

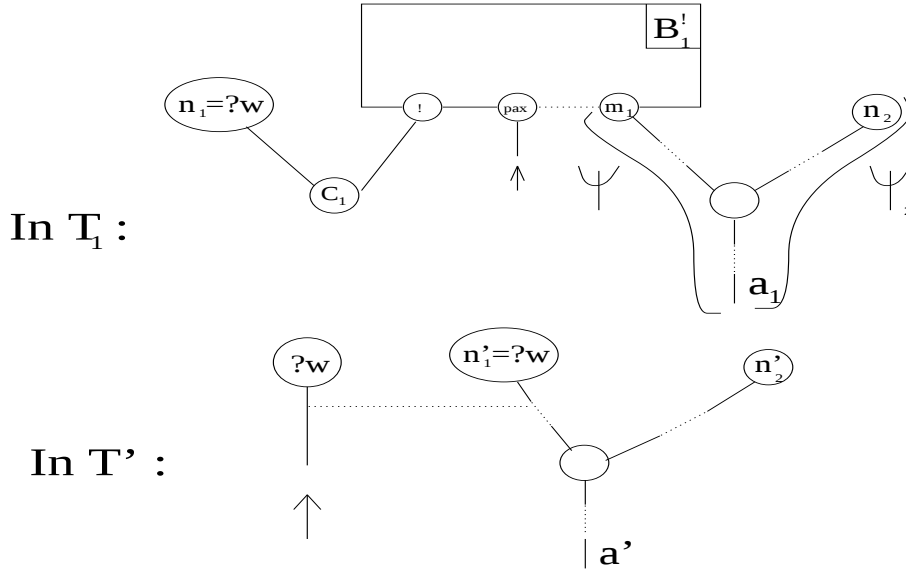


Figure 14.

Case (iii) is represented in figure 15. Let m_1 be the axiom link erased by $[x] = [c_1]$. The existence of the direct paths Ψ'_1 and Φ' of T' gives the existence of a direct path Ψ_1 of T_1 with starting link m_1 and terminal edge a_1 . By hypothesis on σ c_1 is not a residue of c , and then by induction hypothesis c_1 is not a QS -link of T_1 . Because n_1 is a Q -link, c_1 is also a Q -link, so that m_1 is a (main) Q -link which is not an S -link. n_2 is a main S -link which is not a Q -link, and then no direct path with starting link m_1 can cross n_2 . On the other hand, because (2) of remark 2.1.3 holds for n'_2 , it also holds for n_2 and (from lemma A.1.2) there exists a direct path Ψ_2 of T_1 with starting link n_2 and terminal edge a_1 . But in this case a_1 is a QS -edge of T_1 , thus contradicting the induction hypothesis.

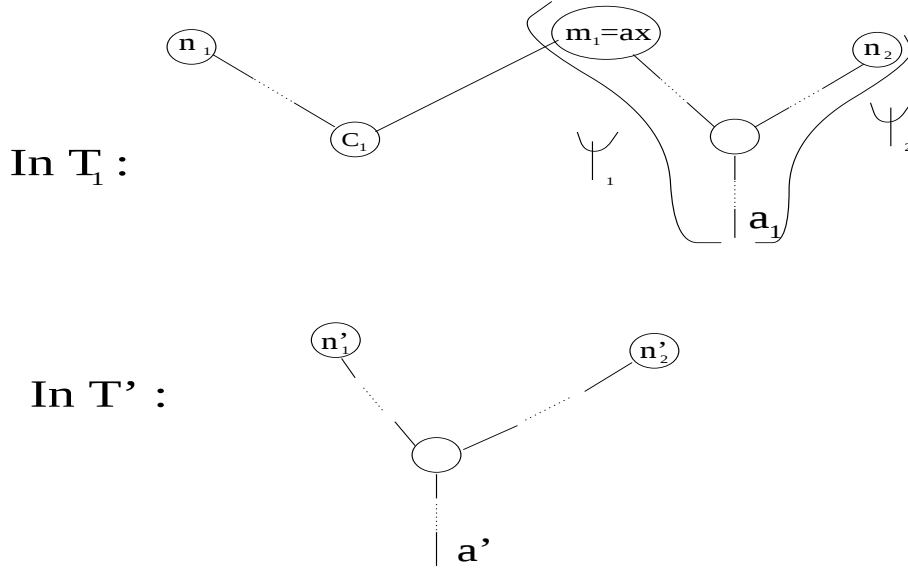


Figure 15.

Part II : We now know that (1) of remark 2.1.3 holds for the link n of T' and we prove the proposition in this case.

Let then c' and Ψ' be the cut link and the direct path associated with n given by remark 2.1.3. We have two possibilities :

(1.1) c' has no lift in T_1 : then x is a logical cut link and c' is created by the e.r.s. $[x]$. Let a'_1 and a'_2 be the two premises of c' in T' , and let l and $l^\perp$ be the dual logical links whose conclusions are the premises of x in T_1 . Because n is a QS -link of T' , either a'_i is a QS -edge (for some $i \in \{1, 2\}$) or $n = c'$ and (for example) a'_1 is a Q -edge while a'_2 is an S -edge. In the first case one among l and $l^\perp$ is a QS -link of T_1 , thus contradicting the induction hypothesis. In the second case x is a QS -link, and then (by induction hypothesis) x is a residue of c in T_1 , thus contradicting the hypothesis on the reduction sequence σ . This means that the case (1.1) cannot occur.

(1.2) c' has the lift c_3 in T_1 : let Φ'_1 (resp. Φ'_2) be the direct path of T' with starting link n'_1 (resp. n'_2) and terminal link c' . For $j = 1, 2$, Φ'_j is obtained by connecting Ψ'_j with Ψ' . Obviously, $n \in \Phi'_1 \cap \Phi'_2$.

Let $j \in \{1, 2\}$. (1) of remark 2.1.3 holds for n_j in T_1 : if (2) holds for n_j , then (by lemma A.1.2) (2) holds for n'_j , which contradicts the fact that (1) holds for n . Let's then denote by c_j and Φ_j the cut link and the direct path of T_1 associated with n_j given by remark 2.1.3. Observe that if $c_1 = c_2$, then

c_j is a QS -link and (by induction hypothesis) c_j is a residue of c in T_1 . This implies that $x \neq c_j$. This means that if $x = c_j$ then $c_1 \neq c_2$.

If c_j has a residue in T' , then (by lemma A.1.2) c' is one of the residues of c_j (i.e. $c_j = c_3$).

If c_j has no residue in T' , then c_j cannot belong to a subproof-net of T_1 erased by $[x]$ (otherwise n_j would have no residue in T'). Then (from remark 2.2.10) $x = c_j$ is a logical cut link (case (i)) or $x = c_j$ is a weakening cut link (case (ii)), or $[x] = [c_j]$ is an axiom e.r.s. (case (iii)).

Case (i) has to be excluded : if $x = c_j$ is a logical cut, then n'_j is the unique residue of n_j and c' is necessarily created by $[x]$. But we are supposing that c' has a lift in T_1 , so that this is impossible.

If case (ii) occurs (for example) for c_1 , then n_1 is the weakening link whose conclusion is a premise of $c_1 = x$, and $c_2 \neq c_1$ has a residue in T' . If case (iii) occurs (for example) for c_1 , then $c_1 \neq c_2$.

This means that we are done, once we have proven the proposition in the three following situations :

- (α) c_1 and c_2 have both c' among their residues : then $c_1 = c_2 = c_3$
- (β) $[x] = [c_1]$ is an axiom e.r.s. : $c_2 \neq c_1$ has a residue in T' , so that $c_2 = c_3$
- (γ) $[x] = [c_1]$ is a $[w]$ e.r.s., and n_1 is the weakening link whose conclusion is a premise of c_1 : then $c_2 = c_3$.

In case (α), $c_3 \in \Phi_1 \cap \Phi_2$ and every link of $\Phi_1 \cap \Phi_2$ is a QS -link, so that by induction hypothesis $\{c_3\} = \Phi_1 \cap \Phi_2$ and c_3 is a residue of c in T_1 . If $[x] \neq [ccad]$, then (by hypothesis on σ) $c_3 \neq x$. This means that we can apply lemma A.1.3 to T_1 and T' : $\{c'\} = \Phi'_1 \cap \Phi'_2$. Then $n = c'$ is a residue of c in T' .

Case (β) is represented in figure 16. Let m_1 be the axiom link erased by $[x] = [c_1]$. Because in T' $c' \in \Phi'_1 \cap \Phi'_2$, there exists a direct path Ψ_1 of T_1 with starting link m_1 and terminal link c_3 . By hypothesis on σ , c_1 is not a residue of c , and then by induction hypothesis c_1 is not a QS -link of T_1 . Because n_1 is a Q -link, c_1 is also a Q -link, so that m_1 is a (main) Q -link which is not an S -link. $c_3 \in \Psi_1 \cap \Phi_2$ and every link of $\Psi_1 \cap \Phi_2$ is a QS -link, so that by induction hypothesis $\{c_3\} = \Psi_1 \cap \Phi_2$ and c_3 is a residue of c in T_1 . This clearly implies that $\{c'\} = \Phi'_1 \cap \Phi'_2$, and then $n = c'$ is a residue of c in T' .

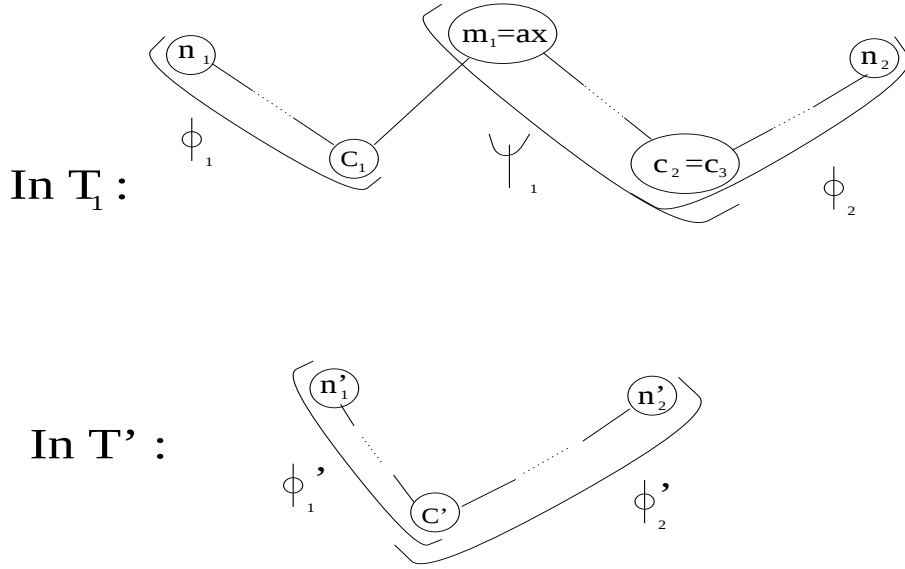


Figure 16.

Case (γ) is represented in figure 17. We argue like in case (ii) of part I. n_1 is a main Q -link and c_1 is then a Q -link. Because $[c_1] \in \sigma$ and $[c_1] \neq [ccad]$, c_1 is not a residue of c in T_1 , and then (by induction hypothesis) c_1 is not an S -link. Let $B_1^\dagger$ be the exponential box erased by $[c_1]$. By induction hypothesis we can apply lemma A.1.6 to $B_1^\dagger$: the conclusion of the main door of $B_1^\dagger$ is a Q -edge (because c_1 is not an S -link) and then every link of $B_1^\dagger$ is a Q -link which is not an S -link.

The existence of Φ_1' in T' implies the existence of a direct path Ψ in T_1 (represented in figure 17) with starting link the pax door m_1 of $B_1^\dagger$ and terminal link c_2 . m_1 is a Q -link while n_2 is a main S -link, so that c_2 is a QS -link and then by induction hypothesis $\{c_2\} = \Psi \cap \Phi_2$ and c_2 is a residue of c in T_1 . Otherwise said the situation is the one in figure 17, and this clearly shows that $\{c'\} = \Phi_1' \cap \Phi_2'$. Then $n = c'$ is a residue of c in T' .

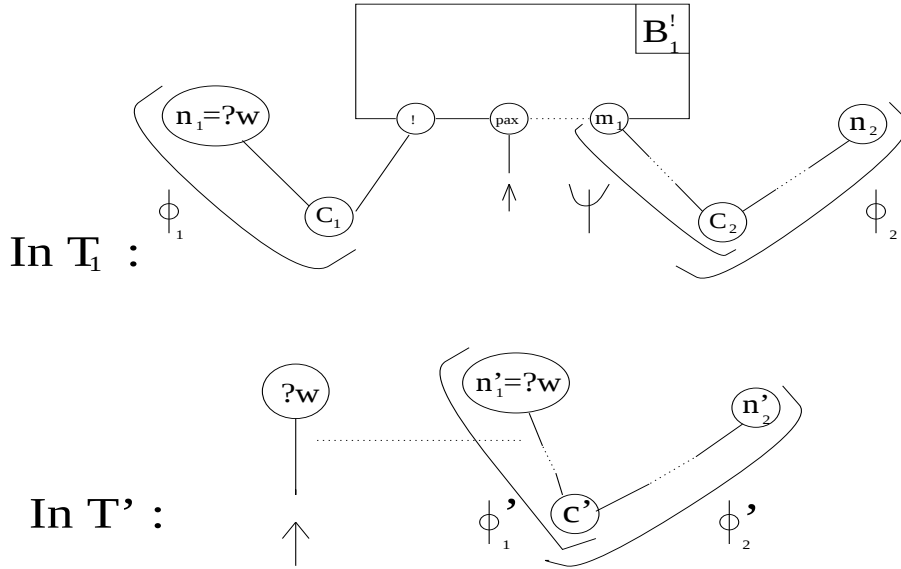


Figure 17.

□

The previous proposition allows to formulate lemma A.1.6 in the following way :

A.1.8. COROLLARY. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c . If $B^!$ is an exponential box of T' , then one (and only one) of the two following conditions holds :

- (1) every link of $B^!$ is a Q -link which is not an S -link
- (2) every link of $B^!$ is an S -link which is not a Q -link.

A.1.9. REMARK. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c . Every residue of c in T' has exponential depth 0 in T' .

This is because c has exponential depth 0 in T , and if $T \xrightarrow{\sigma_1} T_1 \xrightarrow{\sigma_2} T'$ is a decomposition of σ and the residue c_1 of c in T_1 is a cc cut link, then $[c_1] \notin \sigma_2$.

A.1.10. REMARK. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c . Then :

(i) If Φ' is a direct path of T' crossing a Q -link (resp. an S -link) and s.t. the terminal link of Φ' is not a residue of c in T' , then every link of Φ' is a Q -link (resp. an S -link) which is not an S -link (resp. a Q -link). This is an immediate consequence of proposition A.1.7.

(ii) By remark A.1.9, the hypothesis of lemma 3.2.4 are satisfied by the residues of c when $T \xrightarrow{\sigma} T'$, and for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c .

We are now going to prove some lemmas which will be used in section A.4.

A.1.11. LEMMA. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c . Let B' be an additive box of T' and let $i, j \in \{1, 2\}$, $i \neq j$.

If there is a residue c'_i of c in $\gamma_i(B')$, then there is a residue c'_j of c in $\gamma_j(B')$.

proof : By induction on σ . If $T' = T$, then the result is obvious.

Otherwise $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$. Let B_1 be the lift of B' in T_1 , and let c_i be the lift of c'_i in T_1 .

If c_i is a link of B_1 , then it is a link of $\gamma_i(B_1)$, and by induction hypothesis there exists a residue c_j of c in T_1 which is a link of $\gamma_j(B_1)$. If c_j has a residue in T' , then one of these residues is the cut link we are looking for. If c_j has no residue in T' , then $[x] \neq [ccad]$ so that (by hypothesis on σ) $c_j \neq x$. From remark 2.2.10, we deduce that c_j must then belong to a subproof-net of T_1 erased by $[x]$. From remark A.1.9 we deduce that x is a logical $\&/\oplus$ cut link out of the front door of the additive box B_2 of T_1 , and c_j belongs to the component $\gamma_k(B_2)$ of B_2 erased by $[x]$. B_2 is necessarily a box contained in the component j of B_1 (B_2 cannot contain B_1 , otherwise c_i would have no residue in T_1). By induction hypothesis, there exists a residue c_h of c in the component of B_2 different from $\gamma_k(B_2)$. The residue of c_h in T' is the cut link we are looking for.

If c_i is not a link of B_1 , then $[x] = [ccad]$ and, either $[c_i] = [x]$ or c_i is a cut link contained in the subproof-net of T_1 swallowed by B_1 when performing $[x]$. In both cases c_i has two residues : the cut link we are looking for is the residue of c_i different from c'_i . $\square$

A.1.12. LEMMA. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c .

There exists at least a residue of c in T' .

proof : By induction on σ . If $T' = T$, then the result is obvious.

Otherwise $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$. By induction hypothesis, let c_1 be a residue of c in T_1 . If c_1 has a residue in T' , then the result is obvious. Otherwise $[x] \neq [ccad]$, then $x \neq c_1$, so that (again by remark 2.2.10) c_1 belongs to a subproof-net of T_1 erased by $[x]$. From remark A.1.9 we deduce that x is a logical $\&/\oplus$ cut link out of the front door of the additive box B_1 of T_1 , and c_1 belongs to the component of B_1 which is erased by $[x]$. But from lemma A.1.11, the component of B_1 which is not erased by $[x]$ contains a residue c_2 of c in T_1 . The residue of c_2 in T' is a residue of c in T' . $\square$

The notion of jump is crucial to prove the last three lemmas of this section.

A.1.13. LEMMA. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c .

Let $j_{T'}$ be one of the functions associated with σ and given by definition 2.2.8. If m' is a weakening link of $(T', j_{T'})$ which is a (main) Q -link (resp. S -link) then its jumps are all (main) Q -links (resp. S -links).

proof : By induction on σ . If $T' = T$, then the result is obvious.

Otherwise $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$. Let m' be a weakening link of T' which is a (main) Q -link. Let m_1 be the lift of m' in T_1 : m_1 is a (main) Q -link and (by induction hypothesis) its jumps are all Q -links.

If $[x]$ is not an axiom nor a weakening e.r.s., then by (ii) of remark 2.2.9 $j_{T'}(m')$ contains only residues of links of $j_{T_1}(m_1)$. By induction hypothesis the links of $j_{T_1}(m_1)$ are (main) Q -links, and then the links of $j_{T'}(m')$ are all (main) Q -links.

If $[x]$ is a weakening e.r.s., let then n be the of course link whose conclusion is a premise of x , let $B^!$ be the exponential box associated with n and let z_1 be the weakening link whose conclusion is a premise of x . By definition 2.2.8, if for every $z \in j_{T_1}(m_1)$ z is not a link of $B^!$, then $j_{T'}(m')$ is the set of the residues of the links of $j_{T_1}(m_1)$, which are all Q -links. Otherwise by proposition A.1.7 and corollary A.1.8, $B^!$ contains only Q -links that are not S -links, and x, z_1 are Q -links that are not S -links. In this case $j_{T'}(m')$ contains only residues of links of $j_{T_1}(m_1)$ and of $j_{T_1}(z_1)$, then only Q -links. If $[x]$ is an axiom e.r.s., let then n be the axiom link erased by $[x]$. If $n \notin j_{T_1}(m_1)$, then (by definition 2.2.8) $j_{T'}(m')$ is the set of the residues of the links of $j_{T_1}(m_1)$. As usual in this case we conclude by applying the induction

hypothesis. If $n \in j_{T_1}(m_1)$, then we have to prove that every link of $j_{T'}(m')$ which is not the residue of a link of $j_{T_1}(m_1)$ is a main Q -link. Definition 2.2.8 and lemma 2.1.5 show that the lift z_1 in T_1 of every such link z' of T' satisfies one of the two following conditions :

(α) there exists a direct path Ψ_1 of T_1 with starting link z_1 and terminal link x

(β) there exists a weakening link u_1 of T_1 s.t. $z_1 \in j_{T_1}(u_1)$ and there exists a direct path Ψ_1 of T_1 with starting link u_1 and terminal link x .

Observe that $n \in j_{T_1}(m_1)$ is a Q -link (by induction hypothesis), and then x is a Q -link which is not an S -link (by hypothesis on σ). In case (α) (resp. (β)), by remark A.1.10 z_1 (resp. u_1) is a main Q -link. In case (α), this implies that z' is a Q -link and we are done. In case (β), by induction hypothesis we deduce that z_1 is a main Q -link, and then z' is a Q -link. $\square$

A.1.14. LEMMA. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t. $[x] \neq [ccad]$ x is not a residue of c .

Let $j_{T'}$ be one of the functions associated with σ and given by definition 2.2.8, and let $(R', j_{R'})$ be a subproof-net of $(T', j_{T'})$. If n' (resp. m') is a Q -link (resp. an S -link) of T' which is a link of R' , then there is a residue of c in R' .

proof : Let T'_j be the graph obtained from T' by drawing the edges connecting the weakening links of T' and their jumps. By lemma 1.3.14, there exists an oriented path Φ' of T'_j with starting link n' , terminal link m' and crossing only links of R' . Let z' be the first link of Φ' which is an S -link.

If $z' = n'$, n' is a QS -link and then (by proposition A.1.7) a residue of c , so that we are done.

If $z' \neq n'$, let u' be the last link of Φ' which is a Q -link and not an S -link.

Let's denote by $\overrightarrow{u'z'}$ the (oriented) edge of Φ' connecting u' and z' . z' is not a jump of u' : if it were the case, by lemma A.1.13 z' would be a Q -link, then a QS -link and by proposition A.1.7 z' would then be a residue of c in T' , thus contradicting the fact that a jump is an axiom link. For the same reasons, u' is not a jump of z' . This means that there exists an oriented edge α of the proof-net T connecting u' and z' . u' is the source and z' is the goal of α : otherwise u' would be an S -link, thus contradicting the fact that z' is the first S -link crossed by Φ' . Then z' is also a Q -link, which implies (still by proposition A.1.7) that z' is a residue of c in T' . $\square$

A.1.15. LEMMA. Let $T \xrightarrow{\sigma} T'$, and suppose that for every $[x] \in \sigma$ s.t.

$[x] \neq [ccad]$ x is not a residue of c .

If B' is an additive box of T' containing a Q -link and an S -link, then B' contains a residue of c in T' .

proof : By induction on σ . If $T' = T$, then the result is obvious (no additive box of T contains a Q -link and an S -link). Otherwise $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$. Let B_1 be the lift of B' in T_1 .

If B_1 contains an S -link and a Q -link, then by induction hypothesis B_1 contains a residue c_1 of c in T_1 . Suppose for example that c_1 is a link of $\gamma_1(B_1)$. By lemma A.1.11 there exists a residue c_2 of c in T_1 which is a link of $\gamma_2(B_1)$. B_1 has at least a residue in T' (it is B'), so that at least one among c_1 and c_2 has a residue in T' which is a link of B' .

If B_1 contains only (say) Q -links, then one has necessarily $[x] = [ccad]$ and the subproof-net R_1 duplicated by $[x]$ contains some S -links. We have then two possibilities :

- R_1 contains only S -links that are not Q -links : then x is a QS -link of T_1 , that is (by proposition A.1.7) a residue of c in T_1 . In that case any of the two residues of x in T' will do
- R_1 contains S -links and Q -links : then by lemma A.1.14 R_1 contains a residue c_1 of c in T_1 . Any of the two residues of c_1 in T' will then do. $\square$

A.2 The standardization theorem

The following is theorem 4.26.9 p.74 of [Girard 87] :

A.2.1. THEOREM. (Girard 87) All proof-nets of PN_2 are reducible.

If R is any proof-net with conclusions $E_1, \dots, E_k$ and α_i ($i \in \{1, \dots, k\}$) is a *candidat de réductibilité* (see [Girard 87] for the definition) having $E_i^\perp$ among its conclusions, then Girard shows, by induction on the proof-net R , that the proof-net T obtained by cutting R with the α_i is strongly normalizable. The notion of proof-net and the definition of cut-elimination of [Girard 87] are not exactly the ones that we have given in chapter 1 and 2. However, the method of the *candidats de réductibilité* is known to yield a very general proof of strong normalization, and there is not the slightest problem to apply it to the system PN_2 that we have previously defined.

In case R is an additive box $\mathbb{B}$ with front door $D\&C$ and auxiliary doors $A_1, \dots, A_n$, let's call Q_i ($i \in \{1, \dots, n\}$) the *candidat de réductibilité* having $A_i^\perp$

among its conclusions, and S the *candidat de réductibilité* having $D^\perp \oplus C^\perp$ among its conclusions. Girard shows that for S it is enough to take a *candidat de réductibilité* having among its last links a $\oplus$ link with conclusion the edge labeled by $D^\perp \oplus C^\perp$ (which will be one of the two premises of the cut link of T between $\mathbb{B}$ and S). In this paper we will assume that this $\oplus$ link is a $\oplus_1$ link (in case it is a $\oplus_2$ link the proof is of course the same).

Let R_1 (resp. R_2) be the first (resp. the second) component of $\mathbb{B}$, and S_1 be the subproof-net of S obtained by removing the (terminal) $\oplus_1$ link. We call $C_1(T)$ the proof-net obtained by cutting R_1 with the Q_i ($i \in \{1, \dots, n\}$) and with S_1 , and we call $C_2(T)$ the proof-net obtained by cutting R_2 with the Q_i ($i \in \{1, \dots, n\}$).

In Girard's proof of strong normalization, by applying the induction hypothesis one obtains that $C_1(T)$ and $C_2(T)$ are strongly normalizable. What is missing to conclude is then the following

A.2.2. THEOREM. (Additive standardization) If $C_1(T)$ and $C_2(T)$ are strongly normalizable, then so is T .

From now on, T , $C_1(T)$, $C_2(T)$, S and $\mathbb{B}$ will denote proof-nets of the previous form. If $T \xrightarrow{\sigma} T'$, we shall call special cut (resp. special link $\oplus$) any residue of the logical cut between $\mathbb{B}$ and S (resp. any residue of the link $\oplus_1$ of S) in T' .

To prove the theorem, we define a size of every reduct of T , and show that this size shrinks when performing any e.r.s..

A.2.3. REMARK. For the same reasons evoked in remark A.1.1, when $T \xrightarrow{\sigma} T'$, the label of the front door of any residue of $\mathbb{B}$ in T' is $D \& C$, the label of the conclusions of the special $\oplus$ links of T' is $D^\perp \oplus C^\perp$, and the labels of the premises of the special cuts of T' are $D \& C$ and $D^\perp \oplus C^\perp$.

The following useful lemma basically holds because all the residues of $\mathbb{B}$, all the special $\oplus_1$ links and all the special cuts, have exponential depth 0 in T' . The proof of this lemma requires some technicalities, and we suggest to the reader who is not interested in the proof to go directly to definition A.3.1.

A.2.4. LEMMA. Suppose that $T \xrightarrow{\sigma} T'$. Then :

- 1) Let $\&$ (resp. $\oplus_1$) be the front door of a residue of $\mathbb{B}$ (resp. a special $\oplus_1$ link) in T' . There exists a path Φ with starting link $\&$ (resp. $\oplus_1$) and

terminal link a special cut, s.t. every link of Φ different from its starting link and from its terminal link, is a coad link.

- 2) Let c be a special cut in T' , and let Φ be a path having c as starting link and crossing at least one link different from c and from a coad link (such a path obviously exists for any special cut c).

If n is the first link crossed by Φ such that $n \neq c$ and $n \neq \text{coad}$, then n is either a special $\oplus_1$ link either the front door of a residue of $\mathbb{B}$.

Let $T \xrightarrow{\sigma} T'$. Observe that $\mathbb{B}$ and $\oplus_1$ have exponential depth zero in T , so that remark 3.2.1, the independence lemma and lemma A.0.24 apply. The previous lemma shows that every special cut of T' has exponential depth zero in T' , so that also lemma 3.2.4 applies.

A.2.5. DEFINITION. Let Φ be a path in a proof-net, with starting link n and terminal link m . We will say that Φ is a **lucid** path if every link l crossed by Φ and s.t. $l \neq m, n$ is a coad link. We will say that Φ is a **maximal lucid** path if Φ is lucid and $n, m \neq \text{coad}$.

A.2.6. REMARK. (i) If Φ is a lucid path of a proof-net, then either Φ is itself a direct path, or the path Φ^* obtained from Φ by reverting all its edges is a direct path.

(ii) If n is a cut link in a proof-net, then there always exists a maximal lucid path with starting link n .

Lemma A.2.4 is an immediate consequence of the following lemma (just take $R = R_1 = T$, $R' = R_2 = T'$, $\sigma_2 = \sigma$, l the front door of $\mathbb{B}$, $l^\perp$ the unique special $\oplus_1$ link of T and c the unique special cut of T).

A.2.7. LEMMA. Let R be a proof-net s.t. $R \xrightarrow{\sigma_1} R_1 \xrightarrow{\sigma_2} R_2 \xrightarrow{\sigma_3} R'$. Let l and $l^\perp$ be two (dual) propositional logical links of R s.t. the two conclusions of l and $l^\perp$ are the two premises of the logical cut link c , and suppose that l , $l^\perp$ and c have exponential depth zero in R .

If :

- (i.1) For every residue l_1 (resp. $l_1^\perp$) of l (resp. of $l^\perp$) in R_1 , there exists a residue c_1 of c in R_1 and a lucid path Φ_1 with starting link l_1 (resp. $l_1^\perp$) and terminal link c_1
- (ii.1) Every maximal lucid path with starting link a residue c_1 of c in R_1 , has a residue of l or a residue of $l^\perp$ as terminal link

then :

- (i.2) For every residue l_2 (resp. $l_2^\perp$) of l (resp. of $l^\perp$) in R_2 , there exists a residue c_2 of c in R_2 and a lucid path Φ_2 with starting link l_2 (resp. $l_2^\perp$) and terminal link c_2
- (ii.2) Every maximal lucid path with starting link a residue c_2 of c in R_2 , has a residue of l or a residue of $l^\perp$ as terminal link.

proof : By induction on σ_2 . If $R_1 = R_2$, the result is obvious.

Otherwise $R_1 \xrightarrow{\sigma'_2} R'_2 \xrightarrow{[x]} R_2$. In this proof we will constantly use remark 3.2.1 : every residue of l and every residue of c has exponential depth zero.

Let's first prove (i.2). Let l_2 be a residue of l in R_2 and let l'_2 be its lift in R'_2 . By induction hypothesis ((i.2) applied to l'_2), there exists a residue c'_2 of c in R'_2 and a lucid path Φ'_2 with starting link l'_2 and terminal link c'_2 . By induction hypothesis ((ii.2) applied to c'_2), if c'_2 is not a logical cut link, then $[c'_2] = [ccad]$. By remark 2.2.10, the only possibility for c'_2 to have no residue in R_2 would be to disappear in a $[\&/\oplus]$ e.r.s.. But any component of an additive box B of R'_2 containing c'_2 will necessarily contain l'_2 , and c'_2 has a residue in R_2 . The lucid path Φ'_2 of R'_2 is a direct path, so that we can apply lemma A.1.2 : there exists a direct path Φ_2 of R_2 with starting link l_2 and terminal link a residue of c'_2 in R_2 . The proof of lemma A.1.2 shows that if $\Phi_2 \neq \Phi'_2$, then Φ_2 crosses one coad link more or less than Φ'_2 : Φ_2 is lucid.

Let's prove now (ii.2). Let c_2 be a residue of c in R_2 and let Φ_2 be a maximal lucid path of R_2 with starting link c_2 . Let's call l_2 the terminal link of Φ_2 . We have to show that l_2 is a residue of l or a residue of $l^\perp$ in R_2 .

Let m_2 be a main link of R_2 s.t. there exists a direct path Ψ_2 with starting link m_2 , terminal link c_2 , which crosses l_2 and which does not cross any main link of R_2 different from m_2 (we will show that, actually, $m_2 = l_2$). Let m'_2 and c'_2 be, respectively, the lifts in R'_2 of m_2 and c_2 . m'_2 satisfies condition (1) of remark 2.1.3 : otherwise by lemma A.1.2 m_2 would satisfy condition (2) of remark 2.1.3, and we know that this is not the case. Let then Ψ'_2 and c'_3 be the path and the cut link of R'_2 given by remark 2.1.3. We are going to show that $c'_2 = c'_3$. If $c'_2 \neq c'_3$, then by lemma A.1.2 c'_3 has no residue in R_2 . This implies (by remark 2.2.10) that $c'_3 = x$ is a logical cut link (case (i)), or $c'_3 = x$ is a weakening cut link (case (ii)), or $c'_3 = x$ and $[x]$ is an axiom e.r.s. (case (iii)).

In case (i), c_2 is a cut link of R_2 created by $[x]$, with no lift in R'_2 , and this is impossible.

In case (ii), m'_2 is the weakening link whose conclusion is a premise of c'_3 and there exists a direct path Φ'_2 with starting link an auxiliary door of the

biggest exponential box erased by $[x]$ and terminal link c'_2 . By definition of Ψ_2 , Φ'_2 crosses no main link, and this contradicts the induction hypothesis (ii.2) applied to the cut link c'_2 of R'_2 .

In case (iii), there exists a direct path Φ'_2 with starting link the axiom link t'_2 of R'_2 erased by $[x]$ and terminal link c'_2 . By definition of Ψ_2 , t'_2 is the unique main link of Φ'_2 , and this contradicts the induction hypothesis (ii.2) applied to the cut link c'_2 of R'_2 .

We now know that $c'_2 = c'_3$. The induction hypothesis (ii.2) applied to the cut link c'_2 of R'_2 , gives the existence of a maximal lucid path with starting link a residue l'_2 of l or of $l^\perp$ in R'_2 and terminal link c'_2 , which is a subpath of Ψ'_2 . The proof of lemma A.1.2 shows that if $\Psi_2 \neq \Psi'_2$, then Ψ_2 crosses one coad link more or less than Ψ'_2 . This fact, and the fact that Ψ_2 crosses no main link different from m_2 , implies that $l'_2 = m'_2$ and $l_2 = m_2$. Then l_2 is a residue of l or of $l^\perp$ in R_2 . $\square$

A.3 The first projection of the reducts of T

The following is the definition of $C_1(T')$, the analogue of $C_1(T)$ for any reduct T' of T : every residue of $\mathbb{B}$ in T' is replaced by its first component and we erase every special $\oplus_1$ link.

A.3.1. DEFINITION. (first projection) Let T' be such that $T \xrightarrow{\sigma} T'$.

For any residue $\mathbb{B}_{T'}$ of $\mathbb{B}$ in T' with first component $\gamma_1(\mathbb{B}_{T'})$, and for any subproof-net γ of T' with a special link $\oplus_1$ among its last links, we define

$$C_1(T') := T' \left[\gamma_1(\mathbb{B}_{T'}) /_{\mathbb{B}_{T'}} , \gamma \setminus \{\oplus_1\} /_{\gamma} \right]$$

where $\gamma \setminus \{\oplus_1\}$ is obtained from γ by eliminating the special link $\oplus_1$ (then $\gamma \setminus \{\oplus_1\}$ will have $D^\perp$ instead of $D^\perp \oplus C^\perp$ among its conclusions).

As $C_1(T')$ is not exactly defined by substitution (following definition 2.1.6), we still have to define the function $j_{C_1(T')}$. If n is a weakening link of $C_1(T')$, we simply define $j_{C_1(T')}(n)$ as the set of links of $j_{T'}(n)$ which are links of $C_1(T')$.

A.3.2. REMARK. (i) $C_1(T')$ is well-defined thanks to the independence lemma (see ii) of remark 3.2.3) and to lemma A.2.4 : special cuts are replaced by cuts with active formula D . The fact that $(T', j_{T'})$ is a proof-net implies that $(C_1(T'), j_{C_1(T')})$ is a proof-net (condition (1) of definition 1.3.3 for $(T', j_{T'})$ plays of course an essential role).

- (ii) For $T' = T$, $C_1(T)$ is the proof-net previously defined.
- (iii) If in T' there is no residue of $\mathbb{B}$, then from lemma A.2.4 $C_1(T') = T'$
- (iv) Let x be a cut link in the reduct T' of T . If x is not a *logical* special cut, and if x does not appear in the second component of any residue of $\mathbb{B}$ in T' , then we have “the same” cut link x in $C_1(T')$.

The following proposition states that C_1 “projects” a sequence of reductions onto a sequence of reductions/equalities.

A.3.3. PROPOSITION. Let T_1 and T_2 be two proof-nets such that $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T_2$.

One has $C_1(T_1) = C_1(T_2)$ when :

- $[x] = [ccad]$ attractive for a residue of $\mathbb{B}$ in T_1
- x is a logical special cut
- x is a cut link in the second component of a residue of $\mathbb{B}$ in T_1 .

Otherwise $C_1(T_1) \xrightarrow{[x]} C_1(T_2)$.

In particular, if T' is a proof-net such that $T \rightarrow T'$ and if $C_1(T)$ is strongly normalizable, then $C_1(T')$ is also strongly normalizable.

proof : If $[x] = [ccad]$ attractive for a residue of $\mathbb{B}$ in T_1 , x is a logical special cut, or x is a cut link in the second component of a residue of $\mathbb{B}$ in T_1 , then we say that we are in situation (a). Otherwise we say that we are in situation (b).

If there is no residue of $\mathbb{B}$ in T_1 , then there is no residue of $\mathbb{B}$ in T_2 ; we have that $C_1(T_1) = T_1$ and $C_1(T_2) = T_2$, and from lemma A.2.4 there is no special cut in T_1 : we are in situation (b) and we have that $C_1(T_1) \xrightarrow{[x]} C_1(T_2)$.

If x is a logical special cut, then from the definition of C_1 , one clearly deduce $C_1(T_1) = C_1(T_2)$.

If x is a cut link in the first (resp. the second) component of a residue of $\mathbb{B}$ in T_1 , then we are in situation (b) (resp. in situation (a)) and $C_1(T_1) \xrightarrow{[x]} C_1(T_2)$ (resp. $C_1(T_1) = C_1(T_2)$).

If $[x] = [ccad]$ attractive for a residue of $\mathbb{B}$ in T_1 , then the only effect of $[x]$ is to increase the number of links (and maybe of jumps) contained in the residue (without duplicating any residue of $\mathbb{B}$, because of the independence lemma), but the definition of C_1 doesn't care: we are in situation (a) and $C_1(T_1) = C_1(T_2)$.

Otherwise, on the one hand from remark A.3.2 x appears in $C_1(T_1)$, and on the other hand the hypothesis of the substitution lemma (lemma 3.1.1) are satisfied: x is neither logical special neither contained in a residue of

$\mathbb{B}$ in T_1 and, if $[x] = [ccad]$, then $[x]$ is not attractive for any residue of $\mathbb{B}$ in T_1 , and it is impossible for $[x]$ to duplicate a subproof-net of a residue $\mathbb{B}_1$ of $\mathbb{B}$ in T_1 without duplicating the whole $\mathbb{B}_1$. Then, if $\mathbb{B}_1^1, \dots, \mathbb{B}_1^n$ are all the residues of $\mathbb{B}$ in T_1 and $\alpha_1, \dots, \alpha_n$ are any proof-nets s.t. $\mathbb{B}_1^i$ and α_i have the same conclusions, we have, up to substitutions if $[x] = [\forall/\exists]$ (and again thanks to the independence lemma), that $T_1[\alpha_1/\mathbb{B}_1^1, \dots, \alpha_n/\mathbb{B}_1^n] \xrightarrow{[x]} T_2[\alpha_1/\mathbb{B}_1^{1'}, \dots, \alpha_n/\mathbb{B}_1^{n'}]$ where $\mathbb{B}_1^{i'}$ is any residue of $\mathbb{B}_1^i$ in T_2 (there are at most two for each i). In particular, thanks to lemma A.2.4, we can substitute to any residue of $\mathbb{B}$ in T_1 its first component, provided we substitute any special cut by a cut with active formula D : we are in situation b) and $C_1(T_1) \xrightarrow{[x]} C_1(T_2)$. $\square$

Discussion : The first idea would now be to define in a similar way $C_2(T')$ (taking the second component of the residues of $\mathbb{B}$ and erasing the first) for any reduct T' of T , and then to prove the analogue of the previous proposition for C_2 . But here is where the dissymmetry between the role played by the first and the second component of $\mathbb{B}$ appears (in theorem A.2.2 we decided to choose a special link $\oplus_1$ and not $\oplus_2$) : when we reduce a logical special cut, it is always the *second* component of the residue of $\mathbb{B}$ which disappears.

Actually, consider for example the following reduction $T \xrightarrow{[ccad]} T_1 \xrightarrow{[\&/\oplus_1]} T_2$, where :

- T_1 is obtained from T applying an e.r.s. $[ccad]$ attractive for a box B contained in one of the Q_i , in such a way that the subproof-net swallowed by the box B contains $\mathbb{B}$ and the logical special cut,
- T_2 is obtained from T_1 by reducing precisely one of the two logical special cuts (that are residues in T_1 of T 's logical cut between $\mathbb{B}$ and S).

In T_1 , we have two residues of $\mathbb{B}$ as subproof-nets of, respectively, the first and the second component of the residue B_1 of B in T_1 ; but in T_2 , we have only one residue of $\mathbb{B}$, that is a subproof-net of (only) one of the two components of the residue B_2 of B in T_2 .

In T_1 , when we take the second component of the two residues of $\mathbb{B}$, we obtain two proof-nets with C among their conclusions, and we can connect the two edges by a link *coad* (which will be added to the box B_1) to define $C_2(T_1)$. But in T_2 , we have only one residue of $\mathbb{B}$, so that when we take the second component of this residue of $\mathbb{B}$, we obtain a proof-net with C among its conclusions, and there is no way to obtain a proof-net with C among its conclusions from the component of B_2 which does not contain any residue of $\mathbb{B}$. This is because (in the canonical representation of T_1 with respect to

the residue of $\mathbb{B}$ that has a residue in T_2) we have reduced a cut in a *main subproof-net* of T_1 .

The reader is invited to draw a small picture which should clarify the whole thing.

Then, if $T \xrightarrow{\sigma} T'$, we will select a residue $\mathbb{B}'$ of $\mathbb{B}$ in T' , and will distinguish the case in which (taking for any proof-net appearing in σ its canonical representation with respect to the lift of $\mathbb{B}'$ in this proof-net) σ doesn't contain any e.r.s. reducing a cut inside a main subproof-net so that it will be possible to define $C_2^{\mathbb{B}'}(T')$, from the one in which σ contains an e.r.s. reducing a cut inside a main subproof-net, that will be handled in a different way (see proposition A.5.1).

A.4 The second projection of the reducts of T

Let's call Q the subproof-net of T obtained from T by erasing S and the (unique) special cut link of T . We are now in the situation described in section A.1, where the cut c of section A.1 is here the special cut of T between Q and S . We can then use the terminology (and the results) of section A.1.

Let $T \xrightarrow{\sigma} T'$, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' , and let's consider for any proof-net appearing in the reduction sequence σ , its canonical representation with respect to the lift of $\mathbb{B}'$ in this proof-net. The aim of the present section is to show that if σ does not contain any e.r.s. reducing a cut inside a main subproof-net, then one can define $C_2^{\mathbb{B}'}(T')$ by erasing (roughly speaking) all the S -links of T' (see definition A.4.7), and one can prove that $C_2(T) \rightarrow C_2^{\mathbb{B}'}(T')$ (see proposition A.4.9). This is possible, mainly thanks to proposition A.1.7 and lemmas A.1.13 and A.1.15.

The following remark is an immediate consequence of lemma A.2.4.

A.4.1. REMARK. Let $T \xrightarrow{\sigma} T'$. If c' is a special cut of T' , then c' is a QS -link.

The following lemma allows to apply, in the sequel of this section, all the results of section A.1.

A.4.2. LEMMA. Let $T \xrightarrow{\sigma} T'$, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' , and let's consider for any proof-net appearing in the reduction sequence σ , its canonical

representation with respect to the lift of $\mathbb{B}'$ in this proof-net. Suppose that σ does not contain any e.r.s. reducing a cut inside a main subproof-net.

If $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T_2 \xrightarrow{\sigma_2} T'$ is a decomposition of σ and $[x] \neq [ccad]$, then x is not a special cut link.

proof : Let $\mathbb{B}_1$ be the lift of $\mathbb{B}'$ in T_1 . By lemma A.2.4, if x is a special cut link of T_1 and $[x] \neq [ccad]$ then x is a $\&/\oplus$ cut link out of the front door of a residue of $\mathbb{B}$ in T_1 . But this is impossible : on the one hand x cannot be a cut link of a main subproof-net of T_1 (by hypothesis on σ), and on the other hand x cannot be a cut link out of the front door of $\mathbb{B}_1$ (which has a residue in T_2). $\square$

Let $T \xrightarrow{\sigma} T'$, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' , and let's consider for any proof-net appearing in the reduction sequence σ , its canonical representation with respect to the lift of $\mathbb{B}'$ in this proof-net. We are going to prove that in the absence of a reduction inside a main subproof-net, every additive box B of T' which is contained in a main subproof-net and which contains a special cut, “is not new” : as a subproof-net of T' , B is the residue of an additive box containing the lift of $\mathbb{B}'$. This result is an essential ingredient of definition A.4.7 : if B is an additive box of T' containing a special cut, we will have to define $C_2^{\mathbb{B}'}(B)$ only when B contains $\mathbb{B}'$.

In the following lemma, we will speak of “the lift and the residues of the subproof-net $\alpha = B$ (or simply $\alpha = B$)” where B is an additive box, to indicate the fact that we think of B as a *subproof-net* rather than as a box : we then refer to the (stronger) notion of lift and residue of a subproof-net (definition 2.3.7) and not to the (weaker) notion of lift and residue of an additive box (definition 2.2.6).

A.4.3. LEMMA. Let $T \xrightarrow{\sigma} T'$, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' , and let's consider for any proof-net appearing in the reduction sequence σ , its canonical representation with respect to the lift of $\mathbb{B}'$ in this proof-net. Suppose that σ does not contain any e.r.s. reducing a cut inside a main subproof-net.

Let B be an additive box of the main subproof-net β' of T' , suppose that B contains all the special cut links of T' which are links of β' , and that there exists at least one such link. Then, there exists a proof-net R and a decomposition of σ s.t. :

- $T \xrightarrow{\sigma'} R \xrightarrow{\sigma''} T'$, where σ'' is not empty
- there exists a lift $\alpha_R = B_R$ in R of the subproof-net $\alpha = B$ of T' , s.t. $\alpha_R = B_R$ is an additive box containing the lift of $\mathbb{B}'$ in R .

proof : By induction on σ . If $T' = T$, then there is nothing to prove. Otherwise $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T'$.

By hypothesis B contains a special cut link, so that by lemma A.2.4 B contains a residue of $\mathbb{B}$ in T' and B is not itself a residue of $\mathbb{B}$ in T' . Let B_1 be the lift of B in T_1 and $\mathbb{B}_1$ be the lift of $\mathbb{B}'$ in T_1 .

We first show that B_1 contains a residue of $\mathbb{B}$ in T_1 . If no residue of $\mathbb{B}$ in T_1 is contained in B_1 , then $[x]$ is a $[ccad]$ e.r.s. B_1 -attractive, and the subproof-net S_1 of T_1 duplicated by $[x]$ contains a residue of $\mathbb{B}$ in T_1 . By hypothesis on σ , x is not a cut link of one of the main subproof-nets of T_1 , and then S_1 is not a subproof-net of one of the main subproof-nets of T_1 . The only possibility is then that S_1 contains the lift $\mathbb{B}_1$ of $\mathbb{B}'$ in T_1 . But in this case B contains $\mathbb{B}'$, thus contradicting the hypothesis of the lemma.

We now distinguish the two possible cases : B_1 contains or does not contain $\mathbb{B}_1$. First suppose that B_1 contains $\mathbb{B}_1$. Because a residue of B_1 is contained in a main subproof-net of T' , $[x] = [ccad]$ and the subproof-net of T_1 duplicated by $[x]$ contains B_1 . Let α_1 be the subproof-net B_1 of T_1 . α_1 has two residues in T' and $\alpha = B$ is one of them. Then $R = T_1$ is the proof-net that we are looking for, where $\alpha_1 = \alpha_R = B_R = B_1$.

Suppose now that B_1 does not contain $\mathbb{B}_1$. Then B_1 is a box of a main subproof-net β_1 of T_1 . By hypothesis on σ , x is not a link of β_1 and β_1 cannot be erased by $[x]$: β_1 has one or two residue(s) in T' . Let β'_1 be the residue of β_1 containing the subproof-net $\alpha = B$ of T' . $\alpha = B$ has then the subproof-net $\alpha_1 = B_1$ as lift in T_1 . We want to apply the induction hypothesis to the subproof-net $\alpha_1 = B_1$ of T_1 . By lemma A.0.24 β'_1 is a subproof-net of the main subproof-net β' of T' . Then every special cut link c_1 of T_1 which is a link of β_1 has itself a residue c' in T' . Because (by hypothesis on B and β') c' is contained in B , c_1 has to be contained in B_1 . And the lift in T_1 of any special cut of T' contained in B is a link of B_1 . By induction hypothesis, there exists a proof-net R and a decomposition of σ_1 s.t. :

- $T \xrightarrow{\sigma'_1} R \xrightarrow{\sigma''_1} T_1$, where σ''_1 is not empty
- there exists a lift $\alpha_R = B_R$ in R of the subproof-net $\alpha_1 = B_1$ of T_1 , s.t. $\alpha_R = B_R$ is an additive box containing the lift of $\mathbb{B}_1$ in R . Clearly the same proof-net R and the decomposition $T \xrightarrow{\sigma'_1} R \xrightarrow{\sigma''_1} T_1 \xrightarrow{[x]} T'$ of σ will do for the subproof-net $\alpha = B$ of T' . $\square$

The following lemma allows to define $C_2(B)$ when the additive box B of a reduct of T does not contain any S -link (see (ii) of remark A.4.6).

A.4.4. LEMMA. Let $T \xrightarrow{\sigma} T'$ and let B be an additive box of T' containing a residue of $\mathbb{B}$ in T' . Suppose that B does not contain any special cut link and that B is different from a residue of $\mathbb{B}$ in T' .

There exists a unique auxiliary door n of B s.t. the following condition is satisfied :

($\star$) for every residue $\mathbb{B}'$ of $\mathbb{B}$ contained in B , there exists a lucid path Φ' of T' with starting link the front door of $\mathbb{B}'$ and terminal link n .

proof : Let $p_{\mathbb{B}'}$ be the (additive) depth in $\gamma_i(B)$ ($i \in \{1, 2\}$) of the residue $\mathbb{B}'$ of $\mathbb{B}$ in T' . Let $p = \max\{p_{\mathbb{B}'} : \mathbb{B}' \text{ residue of } \mathbb{B} \text{ in } T' \text{ contained in } B\}$. We proceed by induction on p .

If $p = 0$, then let $\mathbb{B}_1$ be a residue of $\mathbb{B}$ contained in B . $\mathbb{B}_1$ has (additive) depth zero in (say) $\gamma_1(B)$. By lemma A.2.4 the conclusion of the front door of $\mathbb{B}_1$ is the premise of a coad link n of B . And still by the same lemma, there exists a lucid path Φ_2 of T' , having as starting link the front door of a residue $\mathbb{B}_2$ of $\mathbb{B}$ in T' which is contained in $\gamma_2(B)$, and as terminal link n . Because $p = 0$, the conclusion of the front door of $\mathbb{B}_2$ is actually the premise of n . The independence lemma implies that there is no residue of $\mathbb{B}$ in T' contained in B and different from $\mathbb{B}_1$ and $\mathbb{B}_2$.

If $p > 0$, there exists a residue $\mathbb{B}_1$ of $\mathbb{B}$ in T' with additive depth greater than zero in (say) $\gamma_1(B)$. Let B_1 be the biggest additive box of $\gamma_1(B)$ containing $\mathbb{B}_1$. $p > 0$ implies that B_1 is different from $\mathbb{B}_1$. The independence lemma implies that all the residues of $\mathbb{B}$ in T' contained $\gamma_1(B)$ are contained in B_1 . By induction hypothesis there exists a unique auxiliary door n_1 of B_1 satisfying condition ($\star$). By lemma A.2.4 and because there is no special cut in B , the conclusion of n_1 is a premise of a coad link n of B . Still by lemma A.2.4 there exists a lucid path Φ_2 having as starting link the front door of a residue $\mathbb{B}_2$ of $\mathbb{B}$ in T' which is contained in $\gamma_2(B)$, and as terminal link n . If $\mathbb{B}_2$ has (additive) depth zero in $\gamma_2(B)$, then the conclusion of its front door is a premise of n and (by the independence lemma) $\mathbb{B}_2$ is the unique residue of $\mathbb{B}$ in T' which is contained in $\gamma_2(B)$: in this case we are done. If $\mathbb{B}_2$ has (additive) depth greater than zero in $\gamma_2(B)$, then we call B_2 the additive box of $\gamma_2(B)$ with the same properties as B_1 in $\gamma_1(B)$. Let n_2 be the coad link of B_2 satisfying condition ($\star$) whose conclusion is a premise of a coad link n' of B . Lemma A.2.4 implies that $n = n'$. $\square$

A.4.5. DEFINITION. Let $T \xrightarrow{\sigma} T'$, and let B be an additive box of T' containing a residue of $\mathbb{B}$ in T' , and s.t. B does not contain any S -link. We define $C_2(B) = B[\gamma_2(\mathbb{B}')/\mathbb{B}']$, where $\mathbb{B}'$ is any residue of $\mathbb{B}$ in T' which is contained in B . Like in the case of C_1 (definition A.3.1), for every weakening link n

of $C_2(B)$, we define $j_{C_2(B)}(n)$ as the set of links of $j_B(n)$ which are links of $C_2(B)$.

In particular, for every residue $\mathbb{B}'$ of $\mathbb{B}$ in T' , $C_2(\mathbb{B}') = \gamma_2(\mathbb{B}')$.

A.4.6. REMARK. Let $T \xrightarrow{\sigma} T'$.

(i) The independence lemma and lemma A.2.4 imply that no special cut can be contained in a residue of $\mathbb{B}$ in T' . Lemma A.1.15 allows then to claim that there is no S -link in any residue of $\mathbb{B}$ in T' (we use this property in the previous definition).

(ii) Let B be an additive box of T' containing a residue of $\mathbb{B}$ in T' , and s.t. B does not contain any S -link. $C_2(B)$ is obviously a proof-net when B is itself a residue of $\mathbb{B}$ in T' . Otherwise, this is still the case because (by remark A.4.1) B does not contain any special cut link, so that we can apply lemma A.4.4 (and the independence lemma) : we substitute the coad link n with premises and conclusion labeled by $D\&C$, by a coad link with premises and conclusion labeled by C . (Of course, like for C_1 , we should mention the fact that $(C_2(B), j_{C_2(B)})$ is a proof-net because (B, j_B) is a proof-net).

The conclusions of $C_2(B)$ are then exactly the conclusions of B , except the edge labeled by $D\&C$ which is replaced by an edge labeled by C .

A.4.7. DEFINITION. (second projection) Let $T \xrightarrow{\sigma} T'$, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' , and let's consider for any proof-net appearing in the reduction sequence σ , its canonical representation with respect to the lift of $\mathbb{B}'$ in this proof-net. Suppose that σ does not contain any e.r.s. reducing a cut inside a main subproof-net. We are going to define $C_2^{\mathbb{B}'}(T')$. We proceed by induction on σ .

If $T' = T$, then we define $C_2^{\mathbb{B}'}(T') = C_2(T)$. Otherwise we have to distinguish among different cases.

Case 1 : If there is no additive box of T' containing $\mathbb{B}'$ and some S -links, then by lemma A.4.4, proposition A.1.7, and lemma A.2.4, T' is the proof-net of figure 18, where B is the biggest box of T' containing $\mathbb{B}'$ (in case $B \neq \mathbb{B}'$), and q (resp. s) is a subgraph of T' containing only Q -links (resp. S -links) that are not S -links (resp. Q -links). In case $B = \mathbb{B}'$ just replace the box B by the box $\mathbb{B}'$, the coad link with conclusion $D\&C$ by a with link with the same conclusion, and erase the residue of $\mathbb{B}$ contained in B . Lemma A.1.13 guarantees that every jump of a weakening link of q is a Q -link, and belongs then to q or to B . The graph R obtained from T' by erasing s and c' is then a proof-net, and we can define $C_2^{\mathbb{B}'}(T') = R[C_2(B)/B]$. $C_2^{\mathbb{B}'}(T')$ is a

proof-net whose conclusions are the Q -edges conclusions of T' and an edge labeled by C .

case 2 : If there exists an additive box of T' containing $\mathbb{B}'$ and some S -links, then let B be one of these boxes, and suppose that $\mathbb{B}'$ is a subproof-net of $\gamma_i(B)$. By lemma A.1.15 and lemma A.1.11, we know that there is a special cut in each of the two components of B . We define $C_2^{\mathbb{B}'}(B)$ and $C_2^{\mathbb{B}'}(\gamma_i(B))$ by induction on $(\sigma$ and on) the maximal (additive) depth d of the special cuts in $\gamma_i(B)$, and we show that the following condition holds :

($\bullet$) the conclusions of $C_2^{\mathbb{B}'}(\gamma_i(B))$ (resp. of $C_2^{\mathbb{B}'}(B)$) are the conclusions of $\gamma_i(B)$ (resp. of B) that are Q -edges and an edge labeled by C .

We first define $C_2^{\mathbb{B}'}(\gamma_i(B))$ in case $d = 0$ (**step 2.1**). Then we define $C_2^{\mathbb{B}'}(\gamma_i(B))$ in case $d > 0$ assuming that we have defined $C_2^{\mathbb{B}'}(\gamma_i(B))$ and $C_2^{\mathbb{B}'}(B)$ when the maximal additive depth d' of the special cuts contained in $\gamma_i(B)$ satisfies $d' < d$ (**step 2.2**). Finally, we define $C_2^{\mathbb{B}'}(B)$ in case $d = 0$ and in case $d > 0$ (the definition is the same in both cases), assuming that we have defined $C_2^{\mathbb{B}'}(\gamma_i(B))$ (**step 2.3**).

step 2.1 : If $d = 0$, then by lemma 3.2.4 there is a unique special cut c_i in $\gamma_i(B)$ and we can argue for $\gamma_i(B)$ exactly like we did before for T' . $\gamma_i(B)$ is the proof-net of figure 9, where B_i is the biggest box of T' which does not contain any S -link and containing $\mathbb{B}'$, and q_i (resp. s_i) is a subgraph of T' containing only Q -links (resp. S -links) that are not S -links (resp. Q -links). Let's call R_i the subproof-net of T' obtained by erasing s_i and c_i . We define $C_2^{\mathbb{B}'}(\gamma_i(B)) = R_i[C_2(B_i)/B_i]$. $C_2^{\mathbb{B}'}(\gamma_i(B))$ satisfies ($\bullet$).

step 2.2 : If $d > 0$, then let B_i be the biggest box of $\gamma_i(B)$ containing $\mathbb{B}'$. By lemma 3.2.4 and lemma A.2.4, B_i contains also all the special cut links of $\gamma_i(B)$. By proposition A.1.7, $\gamma_i(B)$ is the proof-net of figure 10, where the conclusion of the front door of B_i is an edge of q_i (resp. of s_i) when B_i is a Q -box (resp. an S -box), and q_i (resp. s_i) is a subgraph of T' containing only Q -links (resp. S -links) that are not S -links (resp. Q -links). Let α be the component of B_i containing $\mathbb{B}'$. B_i is an additive box of T' containing $\mathbb{B}'$ and some S -links (because B_i contains some special cuts) and s.t. the maximal depth d' of the special cuts in α satisfies $d' < d$. We can then apply the induction hypothesis (on the depth) to B_i : $C_2^{\mathbb{B}'}(B_i)$ is a proof-net satisfying ($\bullet$). Lemma A.1.13 guarantees that every jump of a weakening link of q_i is a Q -link, and belongs then to q_i or to B . The graph R_i obtained from $\gamma_i(B)$ by erasing s_i is then a proof-net, and we can define $C_2^{\mathbb{B}'}(\gamma_i(B)) = R_i[C_2^{\mathbb{B}'}(B_i)/B_i]$, which satisfies ($\bullet$).

step 2.3 : We now define $C_2^{\mathbb{B}'}(B)$.

case 2.3.1 : If B is an S -box, we define $C_2^{\mathbb{B}'}(B) = C_2^{\mathbb{B}'}(\gamma_i(B))$, which satisfies ($\bullet$).

case 2.3.2 : If B is a Q -box, let $\gamma_j(B)$ be the component of B which does not contain $\mathbb{B}'$ (and which is then a main subproof-net of T'). We have to define also $C_2^{\mathbb{B}'}(\gamma_j(B))$. We know (by lemma A.1.11) that there exists a special cut c_j in $\gamma_j(B)$. By lemma 3.2.4, either c_j has additive depth 0 in $\gamma_j(B)$ and it is the unique special cut of $\gamma_j(B)$ (**case 2.3.2.1 :**), or there exists an additive box B_j of $\gamma_j(B)$ containing all the special cuts of T' that are links of $\gamma_j(B)$ (**case 2.3.2.2 :**).

case 2.3.2.1 : In this case, $\gamma_j(B)$ is the proof-net in figure 18, where B_j does not contain any S -link (by lemma A.1.15), and q_j (resp. s_j) is a subgraph of T' containing only Q -links (resp. S -links) that are not S -links (resp. Q -links). For the usual reasons, the graph R_j obtained from $\gamma_j(B)$ by erasing s_j is a proof-net, and we can define $C_2^{\mathbb{B}'}(\gamma_j(B)) = R_j[C_2(B_j)/B_j]$, which satisfies $(\bullet)$.

case 2.3.2.2 : In this case, $\gamma_j(B)$ is the proof-net in figure 19, where B_j is the biggest additive box of $\gamma_j(B)$ containing all the special cuts of T' that are links of $\gamma_j(B)$. The conclusion of the front door of B_j is an edge of q_j (resp. of s_j) when B_j is a Q -box (resp. an S -box), and q_j (resp. s_j) is a subgraph of T' containing only Q -links (resp. S -links) that are not S -links (resp. Q -links). By lemma A.4.3 applied to the subproof-net $\alpha_j = B_j$ of T' , there exists a proof-net R and a decomposition of σ s.t. $T \xrightarrow{\sigma'} R \xrightarrow{\sigma''} T'$, where σ'' is not empty, and the lift $\alpha_R = B_R$ in R of the subproof-net $\alpha_j = B_j$ of T' is an additive box containing the lift $\mathbb{B}_R$ of $\mathbb{B}'$ in R . By induction hypothesis on σ (this is the only place where we use this hypothesis, but it is crucial) $C_2^{\mathbb{B}_R}(B_R)$ is a proof-net satisfying $(\bullet)$. Following definition 2.3.7 B_j and B_R have the same conclusions (up to substitutions). As usual, let R_j be the subproof-net of T' obtained from $\gamma_j(B)$ by erasing s_j . We define $C_2^{\mathbb{B}'}(\gamma_j(B)) = R_j[C_2^{\mathbb{B}_R}(B_R)' / B_j]$, where the “ ’ ” indicates that we have performed in $C_2^{\mathbb{B}_R}(B_R)$ the substitutions required by σ'' .

In both cases (1 and 2), let n be the front door of B and suppose that the premise a_i (resp. a_j) of n which is conclusion of $\gamma_i(B)$ (resp. of $\gamma_j(B)$) is labeled by A_i (resp. A_j) and that the other conclusions of B that are Q -edges are labeled by $E_1, \dots, E_k$. $C_2^{\mathbb{B}'}(\gamma_i(B))$ (resp. $C_2^{\mathbb{B}'}(\gamma_j(B))$) is a proof-net whose conclusions are labeled by $A_i, E_1, \dots, E_k, C$ (resp. $A_j, E_1, \dots, E_k, C$). We can then define $C_2^{\mathbb{B}'}(B)$ as the additive box having as components $C_2^{\mathbb{B}'}(\gamma_i(B))$ and $C_2^{\mathbb{B}'}(\gamma_j(B))$ and satisfying $(\bullet)$.

Finally, let B be the biggest additive box of T' containing $\mathbb{B}'$. We are supposing that B contains some S -links. Then T' is the proof-net in figure 19, where q (resp. s) is a subgraph of T' containing only Q -links (resp. S -links)

that are not S -links (resp. Q -links). As usual, let R be the subproof-net of T' obtained from T' by erasing s . We define $C_2^{\mathbb{B}'}(T') = R[C_2^{\mathbb{B}'}(B)/B]$. The conclusions of $C_2^{\mathbb{B}'}(T')$ are the Q -edges conclusions of T' and an edge labeled by C .

A.4.8. REMARK. Let $T \xrightarrow{\sigma} T'$, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' , and let's consider for any proof-net appearing in the reduction sequence σ , its canonical representation with respect to the lift of $\mathbb{B}'$ in this proof-net. Suppose that σ does not contain any e.r.s. reducing a cut inside a main subproof-net.

Let $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T_2 \xrightarrow{\sigma_2} T'$ be a decomposition of σ , and let $\mathbb{B}_1$ be the lift of $\mathbb{B}'$ in T_1 . If x is a Q -link which is not a link of $\gamma_1(\mathbb{B}_1)$, and if x is not a special cut, then we have “the same” cut link x in $C_2^{\mathbb{B}_1}(T_1)$.

proof : If x is a link of $\mathbb{B}_1$, then the result is obvious.

Otherwise let B_1 be the biggest additive box of T_1 containing $\mathbb{B}_1$. If x is not a link of B_1 , the result is a consequence of definition A.4.7. If x is a link of B_1 , and if B'_1 is the smallest additive box containing x and $\mathbb{B}_1$, then x is contained in the component of B'_1 which contains $\mathbb{B}_1$ (otherwise x would be a cut link contained in a main subproof-net of T_1) and is then (still by definition A.4.7) a link of $C_2^{\mathbb{B}_1}(B'_1)$ and then a link of $C_2^{\mathbb{B}_1}(T_1)$. $\square$

The following proposition is the analogue of proposition A.3.3 : $C_2^{\mathbb{B}'}$ “projects” a sequence of reductions onto a sequence of reductions/equalities.

A.4.9. PROPOSITION. Let $T \xrightarrow{\sigma} T'$, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' , and let's consider for any proof-net appearing in the reduction sequence σ , its canonical representation with respect to the lift of $\mathbb{B}'$ in this proof-net. Suppose that σ does not contain any e.r.s. reducing a cut inside a main subproof-net. Let $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T_2 \xrightarrow{\sigma_2} T'$ be a decomposition of σ and let $\mathbb{B}_1$ (resp. $\mathbb{B}_2$) be the lift of $\mathbb{B}'$ in T_1 (resp. T_2).

One has $C_2^{\mathbb{B}_1}(T_1) = C_2^{\mathbb{B}_2}(T_2)$ when :

- $[x]$ is a $[ccad]$ e.r.s. $\mathbb{B}_1$ -attractive
- $[x]$ is a $[ccad]$ e.r.s. B -attractive, B is a Q -box of T_1 , and the subproof-net of T_1 duplicated by $[x]$ contains only S -links
- x is an S -link
- x is a cut link of $\gamma_1(\mathbb{B}_1)$.

Otherwise, one has $C_2^{\mathbb{B}_1}(T_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(T_2)$.

In particular, if $C_2(T)$ is strongly normalizable, then so is $C_2^{\mathbb{B}'}(T')$.

proof : If $[x]$ is a $[ccad]$ e.r.s. $\mathbb{B}_1$ -attractive, $[x]$ is a $[ccad]$ e.r.s. B -attractive (where B is a Q -box of T_1 , and the subproof-net of T_1 duplicated by $[x]$ contains only S -links), x is an S -link, or x is a cut link of $\gamma_1(\mathbb{B}_1)$, then we say that we are in situation (a). Otherwise we say that we are in situation (b).

If x is a cut link of $\gamma_1(\mathbb{B}_1)$ (resp. of $\gamma_2(\mathbb{B}_1)$), then we are in situation (a) (resp. in situation (b)) and clearly $C_2^{\mathbb{B}_1}(T_1) = C_2^{\mathbb{B}_2}(T_2)$ (resp. $C_2^{\mathbb{B}_1}(T_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(T_2)$). Otherwise, there exists an additive box of T_1 containing $\mathbb{B}_1$ which does not contain x (this box might be $\mathbb{B}_1$ itself). Let B_x be the biggest among these boxes of T_1 . Let B_x^1 be (if it exists) the smallest additive box of T_1 containing B_x and different from B_x . Let $R_1 = T_1$ if B_x^1 does not exist, and let R_1 be the component of B_x^1 which contains B_x otherwise. x is a cut link of R_1 with additive depth zero in R_1 . R_1 is the proof-net of figure 18 or 10, where B_x is denoted by B , q (resp. s) contains only Q -links (resp. S -links) that are not S -links (resp. Q -links).

Observe that R_1 might contain only Q -links that are not S -links : in this case R_1 is the proof-net of figure 19, where s is empty (and $B = B_x$ might be the box $\mathbb{B}_1$). Observe also that when $B_x = \mathbb{B}_1$ and R_1 is the proof-net in figure 18, then we have to replace the box B by the box $\mathbb{B}_1$, the coad link with conclusion $D\&C$ by a with link with the same conclusion, and erase the residue of $\mathbb{B}$ contained in B .

Because x is a cut link with additive depth 0 in R_1 , it is a link of q or a link of s or, in case of figure 18, the special cut c' . We have that $R_1 \xrightarrow{[x]} R_2$, $T_2 = T_1[R_2/R_1]$, and $C_2^{\mathbb{B}_2}(T_2) = C_2^{\mathbb{B}_1}(T_1)[C_2^{\mathbb{B}_2}(R_2)/C_2^{\mathbb{B}_1}(R_1)]$.

Part I : Suppose first that $[x]$ is not a $[ccad]$ e.r.s.. In this case, by lemma A.4.2, x cannot be (in case of figure 18) the special cut c' of T_1 .

If x is a link of s , then by definition A.4.7 $C_2^{\mathbb{B}_1}(R_1) = C_2^{\mathbb{B}_2}(R_2)$: we are in situation (a) and $C_2^{\mathbb{B}_1}(T_1) = C_2^{\mathbb{B}_2}(T_2)$. Observe that in this case, if $[x]$ opens the box B_x (we are necessarily in case of figure 19), then B_x is an S -box of T_1 ; and the previous equalities hold because we know that the component of B_x which is erased by $[x]$ is the one which does not contain $\mathbb{B}_1$: this is precisely the component that we erase when we define $C_2^{\mathbb{B}_1}(B_x)$. To be more precise, suppose that $\mathbb{B}_1$ is contained in the first component of B_x and that α is the residue in T_2 of $\gamma_1(B_x)$. Then we have $C_2^{\mathbb{B}_1}(B_x) = C_2^{\mathbb{B}_1}(\gamma_1(B_x)) = C_2^{\mathbb{B}_2}(\alpha)$.

If x is a link of q , then we are in situation (b). By remark A.4.8 x is also a link of $C_2^{\mathbb{B}_1}(R_1)$. If $[x]$ does not open B_x , then (as a subproof-net) B_x has a unique residue B_x^2 in T_2 and clearly $C_2^{\mathbb{B}_1}(B_x) = C_2^{\mathbb{B}_2}(B_x^2)$. Then $C_2^{\mathbb{B}_1}(R_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(R_2)$, and $C_2^{\mathbb{B}_1}(T_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(T_2)$. If $[x]$ opens the box B_x , then B_x is a

Q -box and $[x]$ opens also the box $C_2^{\mathbb{B}_1}(B_x)$. Definition A.4.7 implies then $C_2^{\mathbb{B}_1}(R_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(R_2)$, and $C_2^{\mathbb{B}_1}(T_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(T_2)$.

Part II : Suppose now that $[x]$ is a $[ccad]$ e.r.s..

If $[x]$ is B_x -attractive, let's call B_x^2 the unique residue of B_x in T_2 . The subproof-net of T_1 duplicated by $[x]$ contains only either Q -links that are not S -links, or S -links that are not Q -links (except the special cut c' when R_1 is the proof-net in figure 18). We then have exactly 3 possibilities :

(1) B_x is an S -box or $B_x = \mathbb{B}_1$: then the only effect of $[x]$ is to increase the number of links of B_x , and because B_x does not exist in $C_2^{\mathbb{B}_1}(T_1)$ this is "invisible" : we are in situation (a) and $C_2^{\mathbb{B}_1}(R_1) = C_2^{\mathbb{B}_2}(R_2)$, and $C_2^{\mathbb{B}_1}(T_1) = C_2^{\mathbb{B}_2}(T_2)$.

(2) B_x is a Q -box and the subproof-net of T_1 duplicated by $[x]$ contains only S -links : then the only effect of $[x]$ is to increase the number of S -links of B_x , and all these links do not exist in $C_2^{\mathbb{B}_1}(T_1)$: we are in situation (a) and $C_2^{\mathbb{B}_1}(B_x) = C_2^{\mathbb{B}_2}(B_x^2)$, $C_2^{\mathbb{B}_1}(R_1) = C_2^{\mathbb{B}_2}(R_2)$, and $C_2^{\mathbb{B}_1}(T_1) = C_2^{\mathbb{B}_2}(T_2)$.

(3) B_x is a Q -box and the subproof-net of T_1 duplicated by $[x]$ contains only Q -links that are not S -links : then the only effect of $[x]$ is to increase the number of links of B_x that are Q -links and not S -links, and all these links do exist in $C_2^{\mathbb{B}_1}(T_1)$: we are in situation (b), $C_2^{\mathbb{B}_1}(R_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(R_2)$, and $C_2^{\mathbb{B}_1}(T_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(T_2)$.

If $[x]$ is not B_x -attractive and $[x]$ does not duplicate B_x , then x is not a special cut link. In this case, it is immediate from definition A.4.7 that when x is a link of q (resp. of s), then we are in situation (b) (resp. in situation (a)) and $C_2^{\mathbb{B}_1}(R_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(R_2)$ (resp. $C_2^{\mathbb{B}_1}(R_1) = C_2^{\mathbb{B}_2}(R_2)$), which implies that $C_2^{\mathbb{B}_1}(T_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(T_2)$ (resp. $C_2^{\mathbb{B}_1}(T_1) = C_2^{\mathbb{B}_2}(T_2)$).

If $[x]$ is B_1 -attractive (for some additive box B_1 of R_1) and $[x]$ duplicates the box B_x , let B_2 be the unique residue of B_1 in T_2 . Let α_1 be the subproof-net (containing B_x) of T_1 duplicated by $[x]$, and let β_1 be the subproof-net of T_1 consisting in the box B_1 , the cut link x and α_1 . α_1 has two residues α_2^1 and α_2^2 in T_2 which are subproof-nets of, respectively, the first and the second component of B_2 . More precisely, for $i \in \{1, 2\}$, $\gamma_i(B_2)$ is obtained by cutting α_2^i and the residue in T_2 of $\gamma_i(B_1)$, by means of the cut link x_2^i (itself a residue of x in T_2). We have to distinguish the case in which B_1 is an S -box (case 1) from the one in which B_1 is a Q -box (case 2).

In case 1 (in this case x might be the special cut c' of figure 18), B_1 is an S -box contained in s which contains no Q -link and x is an S -link. All the Q -links of $\gamma_i(B_2)$ that are not S -links, are then links of α_2^i . If $\mathbb{B}_2$ is a subproof-net of (for example) $\gamma_1(B_2)$, then we have $C_2^{\mathbb{B}_2}(B_2) =$

$C_2^{\mathbb{B}_2}(\gamma_1(B_2)) = C_2^{\mathbb{B}_2}(\alpha_2^1) = C_2^{\mathbb{B}_1}(\alpha_1) = C_2^{\mathbb{B}_1}(\beta_1)$, from which one deduces $C_2^{\mathbb{B}_1}(R_1) = C_2^{\mathbb{B}_2}(R_2)$: we are in situation (a) and $C_2^{\mathbb{B}_1}(T_1) = C_2^{\mathbb{B}_2}(T_2)$.

In case 2, x is a link of q and B_1 is a Q -box contained in q : $C_2^{\mathbb{B}_1}(\gamma_i(B_1)) = \gamma_i(B_1)$. Then $C_2^{\mathbb{B}_2}(\gamma_i(B_2))$ is obtained by cutting $C_2^{\mathbb{B}_1}(\gamma_i(B_1)) = \gamma_i(B_1)$ and $C_2^{\mathbb{B}_1}(\alpha_1)$. This means that $C_2^{\mathbb{B}_2}(B_2)$ can be obtained by applying the “same” e.r.s. $[x]$ to $C_2^{\mathbb{B}_1}(\beta_1)$: $C_2^{\mathbb{B}_1}(\beta_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(B_2)$. Then $C_2^{\mathbb{B}_1}(R_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(R_2)$: we are in situation (b) and $C_2^{\mathbb{B}_1}(T_1) \xrightarrow{[x]} C_2^{\mathbb{B}_2}(T_2)$. $\square$

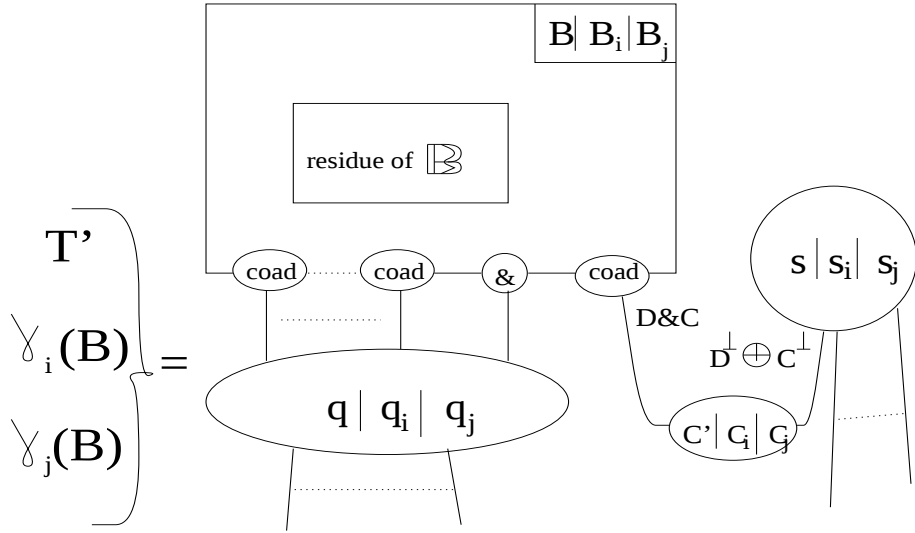


Figure 18.

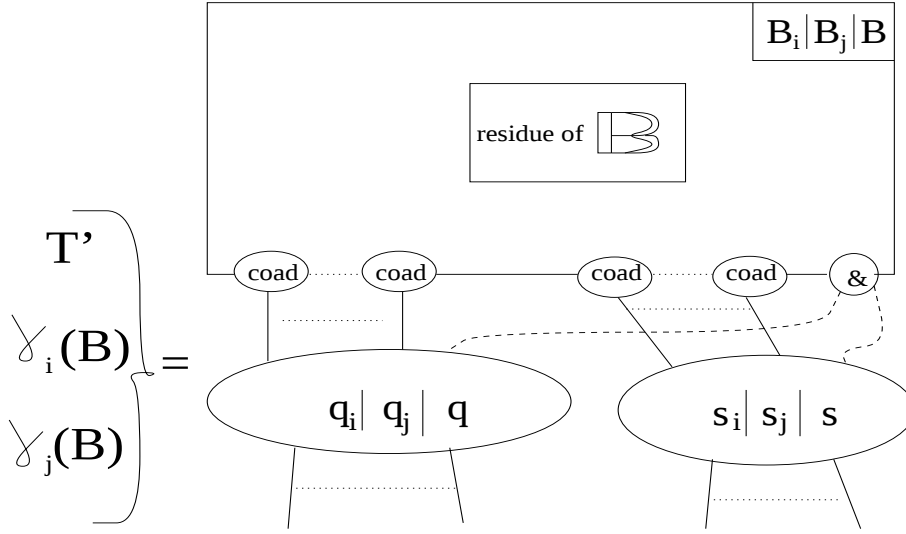


Figure 19.

A.5 Proof of the standardization theorem

A.5.1. PROPOSITION. Let T' be such that $T \xrightarrow{\sigma} T'$.

If $C_2(T)$ is strongly normalizable, then for every residue $\mathbb{B}'$ of $\mathbb{B}$ in T' , $\gamma_2(\mathbb{B}')$ is strongly normalizable.

proof : By induction on σ .

If $T' = T$, then $\gamma_2(\mathbb{B})$ is a subproof-net of $C_2(T)$ which is strongly normalizable by hypothesis.

Otherwise, let $\mathbb{B}'$ be a residue of $\mathbb{B}$ in T' ; we are going to prove that $\gamma_2(\mathbb{B}')$ is strongly normalizable. In the sequel of the proof, all the proof-nets will be considered in their canonical form with respect to the lift of $\mathbb{B}'$ they contain.

If σ doesn't contain any e.r.s. reducing a cut inside a main subproof-net, then from proposition A.4.9, we deduce that $\gamma_2(\mathbb{B}')$ (as a subproof-net of $C_2^{\mathbb{B}'}(T')$) is strongly normalizable.

Otherwise let $[x]$ be the last e.r.s. of σ reducing a cut inside a main subproof-net, more precisely let's consider the following decomposition of σ : $T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T_2 \xrightarrow{\sigma_2} T'$ where σ_2 doesn't contain any e.r.s. reducing a cut inside a main subproof-net, and suppose that $\beta \xrightarrow{[x]} \gamma$ where β (resp. γ) is a main subproof-net of T_1 (resp. of T_2).

From lemma A.0.24, any residue of γ in any of the proof-nets appearing in σ_2 (including T') is a subproof-net of a main subproof-net of this proof-net. But $[x]$ is the last e.r.s. of σ reducing a cut inside a main subproof-net, so that if $[ccad] \in \sigma_2$ this e.r.s. cannot duplicate any subproof-net of a residue of γ without duplicating the whole residue of γ . Then the substitution lemma (lemma 3.1.1) applies : if δ is any proof-net with the same conclusions of γ and if $\gamma'_1, \dots, \gamma'_n$ are the residues of γ in T' , we have that $T_2[\delta/\gamma] \xrightarrow{\sigma_2} T'[\delta/\gamma'_1, \dots, \delta/\gamma'_n]$. In particular, if we take $\delta = \beta$, we have that $T_2[\beta/\gamma] \xrightarrow{\sigma_2} T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n]$ i.e. $T_1 \xrightarrow{\sigma_2} T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n]$. If we define $\sigma' := \sigma \setminus \{[x]\}$, we have that $T \xrightarrow{\sigma'} T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n] \xrightarrow{[x]^n} T'$.

Then, there are two possibilities :

1) there is no residue of γ in T' :

in this case we can eliminate the e.r.s. $[x]$: we have that $T \xrightarrow{\sigma'} T'$, and then, by induction hypothesis, $\gamma_2(\mathbb{B}')$ is strongly normalizable.

2) there is a residue of γ in T' :

by induction hypothesis, if $\mathbb{B}_{T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n]}$ is the lift of $\mathbb{B}'$ in $T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n]$ (w.r.t. the reduction sequence $[x]^n$), then $\gamma_2(\mathbb{B}_{T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n]})$ is strongly normalizable ; and because the n.e.r.s. $[x]$ reduce cuts inside main subproof-nets, we have that $\mathbb{B}_{T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n]} = \mathbb{B}'$ and $\gamma_2(\mathbb{B}_{T'[\beta/\gamma'_1, \dots, \beta/\gamma'_n]}) = \gamma_2(\mathbb{B}')$. $\square$

Let's prove now the (generalized) standardization lemma (theorem A.2.2) : we have to show that T is strongly normalizable.

Consider the following reduction sequence starting from T :

$$T \xrightarrow{\sigma_1} T_1 \xrightarrow{[x]} T_2 \xrightarrow{\sigma_2} T'.$$

For every proof-net R appearing in this reduction sequence and for every residue $\mathbb{B}'$ of $\mathbb{B}$ in R with conclusions labeled by $D\&C$ and $A_1, \dots, A_{n'}$, we define :

- $\mathfrak{R}(R) := \{ \mathbb{B}_R : \mathbb{B}_R \text{ is a residue of } \mathbb{B} \text{ in } R \}$ and $\|\mathfrak{R}(R)\| := \text{card}(\mathfrak{R}(R))$
- $\Lambda_{\mathbb{B}'}(R) := \{ \gamma_1, \dots, \gamma_{n'} : \gamma_i \text{ is the biggest subproof-net of } R \text{ connected to } \mathbb{B}' \text{ through a commutative cut } ccad \text{ having among its premises the conclusion of } \mathbb{B}' \text{ labeled by } A_i \}$. Observe that for some i γ_i may be the empty proof-net.
- $\|\Lambda_{\mathbb{B}'}(R)\| := n(\gamma_1) + \dots + n(\gamma_{n'})$ where $n(\gamma_i)$ = number of links of γ_i .
- $\|\Lambda(R)\| := \sum_{B \in \mathfrak{R}(R)} \|\Lambda_B(R)\|$

We prove that

$$\begin{aligned} & (\|C_1(T_2)\|, \|\mathfrak{R}(T_2)\|, \|\Lambda(T_2)\|, \Sigma_{B \in \mathfrak{R}(T_2)} \|\gamma_2(B)\|) \\ & \qquad \qquad \qquad <_{LEX} \\ & (\|C_1(T_1)\|, \|\mathfrak{R}(T_1)\|, \|\Lambda(T_1)\|, \Sigma_{B \in \mathfrak{R}(T_1)} \|\gamma_2(B)\|) \end{aligned}$$

where $\|C_1(T_i)\|$ (resp. $\|\gamma_2(B)\|$) is the maximal length of the reduction sequences starting from $C_1(T_i)$ (resp. $\gamma_2(B)$).

From proposition A.3.3, if $[x]$ is not a $[ccad]$ e.r.s. which is attractive for a residue of $\mathbb{B}$ in T_1 , if x is not a logical special cut and if x is not a cut in the second component of a residue of $\mathbb{B}$ in T_1 , then $\|C_1(T_2)\| < \|C_1(T_1)\|$.

If x is a logical special cut, then by proposition A.3.3, $\|C_1(T_2)\| = \|C_1(T_1)\|$, and $\|\mathfrak{R}(T_2)\| < \|\mathfrak{R}(T_1)\|$.

If $[x]$ is a $[ccad]$ e.r.s. $\mathbb{B}_1$ -attractive for a residue $\mathbb{B}_1$ of $\mathbb{B}$ in T_1 , then from proposition A.3.3 $\|C_1(T_2)\| = \|C_1(T_1)\|$. Observe that, because of the independence lemma, the number of residues of $\mathbb{B}$ cannot increase in T_2 with respect to T_1 , and we have that $\|\Lambda(T_2)\| < \|\Lambda(T_1)\|$.

If x is a cut in $\gamma_2(\mathbb{B}_1)$ for a residue $\mathbb{B}_1$ of $\mathbb{B}$, then from proposition A.3.3 $\|C_1(T_2)\| = \|C_1(T_1)\|$. We also have $\|\mathfrak{R}(T_2)\| = \|\mathfrak{R}(T_1)\|$, while the independence lemma gives $\|\Lambda(T_2)\| = \|\Lambda(T_1)\|$. If $\mathbb{B}'_1$ is the (unique) residue of $\mathbb{B}_1$ in T_2 , then $\|\gamma_2(\mathbb{B}'_1)\| < \|\gamma_2(\mathbb{B}_1)\|$, so that $\Sigma_{B \in \mathfrak{R}(T_2)} \|\gamma_2(B)\| < \Sigma_{B \in \mathfrak{R}(T_1)} \|\gamma_2(B)\|$.

□

A.5.2. REMARK. Observe that to have a complete proof of strong normalization for this revisited PN_2 , one still has to prove that the standardization lemma for the multiplicative and exponential fragment proven in [Danos 90] still holds in our framework.

Annexe B

Denotational semantics for polarized (but non-constrained) LK by means of the additives

B.1 Introduction

Much work has been done in the past 6-7 years to extract the computational content from classical proofs. On the proof-theoretical side, let's quote for example [Girard 91b] (LC), [Parigot 91] (FD), [DJS 97] (LK^{tq} and its restrictions) and [BarBer 95] (the symmetric λ -calculus).

With the notable exception of the symmetric λ -calculus of [BarBer 95] which enjoys only strong normalization, all the systems previously mentioned enjoy the usual good computational properties (strong normalization and confluence) and have a denotational semantics. To obtain such good properties, it is necessary to cope with the “non-determinism” of classical cut-elimination : this is done by equipping classical formulas with an index, a colour or a polarity (depending on the personal taste of the authors), which will tell us how to perform cut-elimination in case of conflict.

But it turned out (perhaps a bit surprisingly) that adding this information is not enough to fix classical logic : there are rules that are “difficult” to handle, in the sense that if we don't restrict their application to the premises of a particular form, then the elimination of certain ‘logical cuts’ (Gentzen's key-cases) yields *two* proofs. What is catastrophic is the fact that the identification of these two proofs leads to the collapse of the whole system :

we have to identify all the proofs of a given sequent (see [Girard 91b] and [DJS 97]).

The solution proposed by the works previously mentioned is either to avoid such difficult rules (this is possible keeping completeness w.r.t. classical provability, see [Parigot 91] and [BarBer 95]) or to restrict the application of such difficult rules to the premises which lead to key-cases the elimination of which is unproblematic (see [Girard 91b] and [DJS 97]).

A trace of this restriction is the fact that the translations from LK (with polarized or coloured formulas) to LC or LK^{tq} (and its subsystems) are far from being deterministic : there are two possible ways to translate the difficult rules with certain premises. The same holds for the translations in linear logic (LL), which are the actual starting point of both the works [Girard 91b] and [DJS 97].

However, these translations suggest that the computational object which lies ‘behind’ a classical proof *is not* a linear proof, but rather a superimposition of linear proofs. This work is a first attempt to turn this intuition into a mathematical statement : “there exists a denotational semantics of polarized (but non-constrained) LK in coherent spaces”.

Trying to take seriously the idea of superimposition, we define inductively an embedding P^a of polarized LK -proofs in LL , which is an extension of the P -embedding defined by Girard in [Girard 91b] : when we meet a “difficult” rule, then, instead of choosing one of the two possible linear translations, we superimpose them.

Technically speaking, the obvious candidate to superimpose proofs is the linear connective “&”. Following one of the main points of chapter 3 (namely the “independence lemma”), to make sure that the components of the additive boxes introduced by P^a do not interact during cut-elimination, we keep these additive boxes at exponential depth zero. This is done by choosing as main doors for these additive boxes a “purely negative” formula (in the sense of [Girard 91b]), namely $\perp \& \perp$: this makes the additive boxes “commute” (semantically speaking) with all the exponential ones.

Thus, we translate the classical proof α by the linear proof $P^a(\alpha)$, and then, by “delaying” as much as possible all the additive boxes introduced by P^a , we obtain a superimposition of linear proofs with the same denotational semantics of $P^a(\alpha)$.

This suggests to define a classical cut-elimination step as a cut-elimination step in *one* of the “slices” that the classical proof superimposes. We obtain in such a way a denotational semantics (induced by P^a) for polarized LK , which still satisfies all the computational isomorphisms of LC (see [Girard 91b]).

B.2 Preliminaries

B.2.1 Linear logic and classical logic

The use of LL as a tool to analyze classical cut-elimination is at the core of [Girard 91b] and [DJS 97].

Girard observes that if A and B are linear formulas of type “!” (resp. of type “?”), that is if $A \vdash !A$ (resp. $?A \vdash A$) is provable in LL , then so are the formulas $A \otimes B$ and $A \oplus B$ (resp. $A \wp B$ and $A \& B$), and that A is of type “!” (resp. “?”) iff $A^\perp$ is of type “?” (resp. “!”). This suggests a translation P from classical to linear logic, which translates each classical formula by a linear formula of type “!” or of type “?”. To apply structural operations directly to formulas of type “?” (and not only to formulas starting with “?”), Girard slightly modifies the coherent denotational semantics of LL and introduces the semantics of correlation spaces for the linear fragment necessary to interpret classical logic. He also proves that the main properties of classical connectives can be recovered at the computational level (i.e. are isomorphisms in the semantics).

Danos, Joinet and Schellinx work directly with special translations in LL (called “decorations”) which suggest a cut-elimination procedure (the tq -protocol) for different versions of coloured second order classical sequent calculus (LK^{tq} , LK^η , LK_{pol}^η). They show that Girard’s LC is a subsystem of LK_{pol}^η (up to the stoup/no stoup formulation, see subsection 2.3), that the P -embedding defined by Girard is suitable for the whole (first order) LK_{pol}^η , and that the cut-elimination procedure defined by Girard for LC is the tq -protocol restricted to LC .

In [QuatrTorto 96], we prove that the P -embedding induces a denotational semantics (to which we will also refer to as the P -semantics) in coherent spaces for first order LK_{pol}^η (then for LC): if π and π' are LK_{pol}^η -proofs, from $\pi \rightarrow_{tq} \pi'$ we deduce $[P(\pi)] = [P(\pi')]$ (where $[d]$ denotes the coherent semantics of the linear proof d). In section 15.5, we clarify the links between first order LK_{pol}^η and LC : LC is obtained from the multiplicative fragment of LK_{pol}^η by decomposing the multiplicative conjunction between two negative formulas into an additive one and a multiplicative conjunction with “ V ” (the classical unit for the multiplicative conjunction).

B.2.2 About the additives

In this subsection, we briefly recall the independence lemma, proven in chapter 3. Although we will not use this result (so that this subsection is not,

strictly speaking, necessary to understand the content of the paper), it has been one of the main sources of inspiration for the present work. We assume here that the reader has some acquaintance with the proof-nets of LL defined in [Girard 87].

If T and T' are two proof-nets s.t. $T \rightarrow T'$ (i.e. T' is obtained from T by applying some steps of cut-elimination), then any additive box B of T' comes from exactly one additive box $\overleftarrow{B}$ of T : we say that B is a residue of $\overleftarrow{B}$ in T' .

B.2.3. LEMMA. (Independence lemma). Let T, T' be two proof-nets, s.t. $T \rightarrow T'$. Let $\mathbb{B}$ be an additive box with exponential depth 0 in T , and $\mathbb{B}'$ a residue of $\mathbb{B}$ in T' . If $\mathbb{B}'$ has additive depth p and $B_1, \dots, B_p$ are the additive boxes containing $\mathbb{B}'$ in T' , then for all residue $\mathbb{B}''$ of $\mathbb{B}$ in T' (different from $\mathbb{B}'$), there exists $i \in \{1, \dots, p\}$ s.t. $\mathbb{B}''$ is a subproof-net of the component of B_i which does not contain $\mathbb{B}'$.

A possible way of reading this lemma (connected to the present paper) is the following : “normalization preserves superimposition at exponential depth 0”.

B.2.4 Notations and conventions

We shall use the conventions of [Girard 91b], thus distinguish between *positive* and *negative* formulas and take a one-sided version of sequent calculus, rather than distinguish between formulas coloured q and t and take a two-sided version of sequent calculus like in [DJS 97].

For classical formulas we take the following conventions :

$P, Q, R, \dots$ will denote positive formulas

$L, M, N, \dots$ will denote negative formulas

$A, B, C, \dots$ will denote formulas that may be either positive or negative.

We say that a formula (or, better said, an occurrence of formula) A in a proof π of any subsystem of LK is *main*, if A is either one of the (two) conclusions of an axiom, or the main conclusion of a logical rule of π .

In the sequel we will call LC the extension of Girard’s original system to first order LK_{pol}^η . (We don’t call it LK_{pol}^η because we will use a version which makes use of “the stoup”, and this is more in the spirit of LC , and also because we don’t want to create confusions with the system LK_{pol} defined in the next section).

To be precise one should note that there are LC -rules that are not (strictly speaking) rules of LK_{pol}^η . This is because the constraint on proofs induced by Girard’s stoup is (slightly) less strong than [DJS 97]’s η -constraint : an attractive formula (i.e. necessary positive in the one-sided version) which

is active in an irreversible rule has to be main following [DJS 97], and only main up to reversible, structural and (some kind of) cut rules on the context, following [Girard 91b]. From the semantical point of view the two constraints are the same : to each LC -proof one can associate several LK_{pol}^η -proofs, but the denotational semantics induced by the P -embedding in LL equates all of them. But syntactically speaking, in LC there are for example two ways of performing certain logical steps : if $\gamma = cut(\alpha, \beta)$ is a proof of $\vdash \Gamma, \Delta; \Pi$ obtained from the proof α of $\vdash \Gamma; R \wedge_m Q$ and the proof β of $\vdash \Delta, \neg R \vee_m \neg Q; \Pi$ by applying a logical cut rule, and if we call γ_1 and γ_2 the two proofs obtained from γ by performing the two cuts with active formulas $R/\neg R, Q/\neg Q$ in the two possible ways, then γ_1 and γ_2 might have two syntactically different normal forms w.r.t. the tq -protocol of cut-elimination (of course the normal forms will anyway be equal in the semantics).

So, what we call LC in this paper, is the system LK_{pol}^η (with the stronger η -constraint on proofs), the only *notational* difference being that we will write sequents making use of the stoup (this is particularly helpful for our purposes, as it will appear clearly in the next section).

In the whole paper, we will apply directly the structural operations to any linear formula A of type “?” , implicitly meaning (if A does not actually start with “?”) that we first apply a dereliction to A , then the structural operation to $?A$, and finally we cut with the proof of $?A \vdash A$ (with $\vdash A, !A^\perp$ in the one-sided version). One can also think directly in terms of correlation spaces (which, by the way, is proven to be the same in [QuatrTorto 96]).

If B and C are the immediate subformulas of A , an η -proof of A is the proof of $\vdash \neg A, A$ consisting in the axioms $\vdash \neg B, B$ and $\vdash \neg C, C$, and precisely one instance of each of the logical rules introducing A (or $\neg A$)’s main connective. We will also call “piece(s) of η -proof” the subproof(s) of η , obtained from η by removing the last (reversible) rule.

B.3 LK_{pol} and the embedding P^a in linear logic

We consider here a version of the sequent calculus LK where each formula comes equipped with a polarity, in the style of [Girard 91b].

The real difference between the approach we present here and the ones previously mentioned lies in the fact that we *do not* restrict the rules of the calculus LK . Both in [Girard 91b] and in [DJS 97] there is a restriction in the application of certain rules (“the irreversible ones in which one of the active formulas of one of the premises is positive”) : as stated, this is the reason for the introduction of “the stoup” in [Girard 91b] and of the

“ η -constraint” in [DJS 97].

B.3.1 The sequent calculus LK_{pol}

LK_{pol} is nothing but the usual calculus LK (where we distinguish between additive and multiplicative rules) for first order classical logic. However, in LK_{pol} , each formula comes equipped with a polarity (it is positive or negative), which is completely determined by the polarity of the atoms : the polarity of compound formulas is obtained by looking at the main connective of the formula, if it is $\vee_m$ or $\wedge_a$ (resp. $\wedge_m$ or $\vee_a$) then the formula is negative (resp. positive). Of course, for any formula A , A and $\neg A$ have different polarities (see [Girard 91b] and [DJS 97] for more details).

Moreover, we “decorate” (not in the technical sense of [DJS 97]!) each sequent of this polarized LK by means of the symbol “;”, where on the right of “;” there is at most one *positive* formula. This is of course in the spirit of [Girard 91b]. However, our calculus does not put any constraint on proofs. The following system (which will be called LK_{pol}) shows how it is possible to “decorate” by means of “;”, in a unique way, *any* polarized proof α of LK , in such a way that all the sequents of α have the form $\vdash \Gamma; \Pi$, where Γ is a set of (polarized) formulas and Π is a set containing at most one formula, which has to be positive.

Axiom :

$$\vdash \neg P; P$$

Cut-rules :

$$\frac{\begin{array}{c} \gamma \\ \vdots \\ \vdash \Gamma; P \end{array} \quad \begin{array}{c} \delta \\ \vdots \\ \vdash \Delta, \neg P; \Pi \end{array}}{\vdash \Gamma, \Delta; \Pi} \quad \frac{\begin{array}{c} \gamma \\ \vdots \\ \vdash \Gamma, P; \Pi \end{array} \quad \begin{array}{c} \delta \\ \vdots \\ \vdash \Delta, \neg P; \Pi' \end{array}}{\vdash \Gamma, \Delta, \Pi'; \Pi}$$

Structural rules :

$$\frac{\begin{array}{c} \beta \\ \vdots \\ \vdash \Gamma; \Pi \end{array}}{\vdash \Gamma, A; \Pi} \quad \frac{\begin{array}{c} \beta \\ \vdots \\ \vdash \Gamma, A, A; \Pi \end{array}}{\vdash \Gamma, A; \Pi}$$

Logical rules :

1) constants :

$$\vdash; V \quad \vdash \neg F; \Pi$$

2) unary logical rules :

$$\begin{array}{c}
 \beta \\
 \vdots \\
 \frac{\vdash \Gamma, A, B; \Pi}{\vdash \Gamma, A \vee_m B; \Pi} \\
 \beta
 \end{array}
 \quad
 \begin{array}{c}
 \beta \\
 \vdots \\
 \frac{\vdash \Gamma, A; P}{\vdash \Gamma, A \vee_m P;} \\
 \beta
 \end{array}
 \quad
 \begin{array}{c}
 \beta \\
 \vdots \\
 \frac{\vdash \Gamma; P}{\vdash \Gamma; P \vee_a B} \\
 \beta
 \end{array}
 \quad
 \begin{array}{c}
 \beta \\
 \vdots \\
 \frac{\vdash \Gamma, P; \Pi}{\vdash \Gamma, P \vee_a B; \Pi} \\
 \beta
 \end{array}
 \quad
 \begin{array}{c}
 \beta \\
 \vdots \\
 \frac{\vdash \Gamma, N; \Pi}{\vdash \Gamma, \Pi; N \vee_a B} \\
 \beta
 \end{array}$$

3) binary logical rules :

For the connective $\wedge_a$ we have

$$\begin{array}{c}
 \gamma \quad \delta \\
 \vdots \quad \vdots \\
 \frac{\vdash \Gamma, A, \Pi'; \Pi \quad \vdash \Gamma, B, \Pi; \Pi'}{\vdash \Gamma, A \wedge_a B, \Pi, \Pi';} \\
 \gamma \quad \delta \\
 \vdots \quad \vdots \\
 \frac{\vdash \Gamma, A; \Pi \quad \vdash \Gamma, B; \Pi}{\vdash \Gamma, A \wedge_a B; \Pi}
 \end{array}$$

where we obviously suppose $\Pi \neq \Pi'$, and

$$\begin{array}{c}
 \gamma \quad \delta \\
 \vdots \quad \vdots \\
 \frac{\vdash \Gamma, A; \Pi \quad \vdash \Gamma, \Pi; P}{\vdash \Gamma, A \wedge_a P, \Pi;} \\
 \gamma \quad \delta \\
 \vdots \quad \vdots \\
 \frac{\vdash \Gamma; P \quad \vdash \Gamma; Q}{\vdash \Gamma, P \wedge_a Q;}
 \end{array}$$

The case of the connective $\wedge_m$ is the more complicated, and we have to distinguish all the possible cases for the polarity of the active formulas :

$$\begin{array}{c}
 \gamma \quad \delta \quad \gamma \quad \delta \quad \gamma \quad \delta \\
 \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \\
 \frac{\vdash \Gamma; P \quad \vdash \Delta; Q}{\vdash \Gamma, \Delta; P \wedge_m Q} \quad
 \frac{\vdash \Gamma, P; \Pi \quad \vdash \Delta, Q; \Pi'}{\vdash \Gamma, \Delta, P \wedge_m Q, \Pi, \Pi';} \quad
 \frac{\vdash \Gamma; P \quad \vdash \Delta, Q; \Pi}{\vdash \Gamma, \Delta, P \wedge_m Q; \Pi} \\
 \gamma \quad \delta \quad \gamma \quad \delta \quad \gamma \quad \delta \\
 \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \\
 \frac{\vdash \Gamma, M; \Pi \quad \vdash \Delta, N; \Pi'}{\vdash \Gamma, \Delta, \Pi, \Pi'; M \wedge_m N} \quad
 \frac{\vdash \Gamma, P; \Pi \quad \vdash \Delta, N; \Pi'}{\vdash \Gamma, \Delta, P \wedge_m N, \Pi'; \Pi} \quad
 \frac{\vdash \Gamma; P \quad \vdash \Delta, N; \Pi'}{\vdash \Gamma, \Delta, \Pi'; P \wedge_m N}
 \end{array}$$

First order quantifiers :

$$\begin{array}{c}
 \beta \\
 \vdots \\
 \frac{\vdash \Gamma, A; \Pi}{\vdash \Gamma, \forall x A; \Pi} \\
 \beta
 \end{array}
 \quad
 \begin{array}{c}
 \beta \\
 \vdots \\
 \frac{\vdash \Gamma; P}{\vdash \Gamma, \forall x P;} \\
 \beta
 \end{array}$$

However, even among these rules, there is a rule which has to be carefully translated : it is the $\wedge_a$. We proceed in the following way (if the main conclusion of the rule $\wedge_a$ is $A \wedge B$, suppose for example that A is positive and not in the stoup, B is negative, and the two formulas in the two stoups are different) :

$$\begin{array}{c}
P^a(\gamma) \\
\vdots \\
\frac{\frac{\vdash ?\Gamma, ?P, ?\Pi', \Pi, \vec{\perp}_1}{\vdash ?\Gamma, ?P, ?\Pi', ?\Pi, \vec{\perp}_1}}{\vdash ?\Gamma, ?P, ?\Pi', ?\Pi, \vec{\perp}_1, \vec{\perp}_2} W \\
\vdash ?\Gamma, ?P, ?\Pi', ?\Pi, \vec{\perp}_1, \vec{\perp}_2
\end{array}
\quad
\begin{array}{c}
P^a(\delta) \\
\vdots \\
\frac{\frac{\vdash ?\Gamma, N, ?\Pi, \Pi', \vec{\perp}_2}{\vdash ?\Gamma, N, ?\Pi, ?\Pi', \vec{\perp}_2}}{\vdash ?\Gamma, N, ?\Pi, ?\Pi', \vec{\perp}_1, \vec{\perp}_2} W \\
\vdash ?\Gamma, N, ?\Pi, ?\Pi', \vec{\perp}_1, \vec{\perp}_2
\end{array}$$

The important point here is that we do not “mix” formulas coming from $\vec{\perp}_1$ and formulas coming from $\vec{\perp}_2$, even if they are the same (see also the remark following definition B.4.4).

- for the rules of LK_{pol} that are not rules of LC , let's consider some of the rules introducing the connective $\wedge_m$ which contain all the difficult cases :
 - suppose that in the premisses the active formula P is out of the stoup and the main conclusion of the rule is $P \wedge_m N$, then we translate by

$$\begin{array}{c}
P^a(\delta) \\
\vdots \\
\frac{\frac{\vdash ?\Delta, N, \Pi', \vec{\perp}_2}{\vdash ?\Delta, N, ?\Pi', \vec{\perp}_2}}{\vdash ?\Delta, !N, ?\Pi', \vec{\perp}_2}
\end{array}
\quad
\begin{array}{c}
P^a(\gamma) \\
\vdots \\
\frac{\frac{\vdash ?\Gamma, ?P, \Pi, \vec{\perp}_1}{\vdash ?\Gamma, ?(P \otimes !N), ?N^\perp, \Pi, \vec{\perp}_1}}{\vdash ?\Gamma, ?\Delta, ?(P \otimes !N), ?\Pi', \Pi, \vec{\perp}_1, \vec{\perp}_2}
\end{array}$$

Observe that the following is another way to translate this rule, but the proof thus obtained has the same denotational semantics as the previous one :

$$\begin{array}{c}
P^a(\delta) \\
\vdots \\
\frac{\frac{\frac{\frac{\vdash ?\Delta, N, \Pi', \vec{\perp}_2}{\vdash ?\Delta, N, ?\Pi', \vec{\perp}_2}}{\vdash ?\Delta, !N, ?\Pi', \vec{\perp}_2}}{\vdash ?\Delta, P \otimes !N, P^\perp, ?\Pi', \vec{\perp}_2}}{\vdash ?\Delta, ?(P \otimes !N), P^\perp, ?\Pi', \vec{\perp}_2}}
\end{array}$$

By inspection of cases, one can check that there is a unique (semantically speaking) canonical way to translate all the rules introducing $\wedge_m$, except in case both the active formulas are positive and out of the stoup.

- suppose now that in the premisses both the active formulas P and Q are out of the stoup (this is *the* crucial case), then we translate by :

$$\begin{array}{c}
\begin{array}{c}
P^a(\gamma) \\
\vdots \\
\frac{\vdash ?\Gamma, ?P, \Pi, \perp_1 \quad \frac{\frac{\vdash P, P^\perp \quad \vdash Q, Q^\perp}{\vdash P \otimes Q, P^\perp, Q^\perp}}{\vdash ?(P \otimes Q), P^\perp, Q^\perp}}{\vdash ?\Gamma, ?(P \otimes Q), !P^\perp, Q^\perp} \\
\vdots \\
\frac{\vdash ?\Gamma, ?(P \otimes Q), Q^\perp, \Pi, \perp_1}{\vdash ?\Gamma, ?(P \otimes Q), Q^\perp, ?\Pi, \perp_1} \\
\vdots \\
\frac{\vdash ?\Gamma, ?(P \otimes Q), !Q^\perp, ?\Pi, \perp_1}{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2} \\
\vdots \\
\frac{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2}{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2} \\
\vdots \\
\frac{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2}{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2, \perp}
\end{array}
\quad
\begin{array}{c}
P^a(\delta) \\
\vdots \\
\frac{\vdash ?\Delta, ?Q, \Pi', \perp_2 \quad \frac{\frac{\vdash P, P^\perp \quad \vdash Q, Q^\perp}{\vdash P \otimes Q, P^\perp, Q^\perp}}{\vdash ?(P \otimes Q), P^\perp, Q^\perp}}{\vdash ?\Delta, ?(P \otimes Q), P^\perp, !Q^\perp} \\
\vdots \\
\frac{\vdash ?\Delta, ?(P \otimes Q), P^\perp, \Pi', \perp_2}{\vdash ?\Delta, ?(P \otimes Q), P^\perp, ?\Pi', \perp_2} \\
\vdots \\
\frac{\vdash ?\Delta, ?(P \otimes Q), !P^\perp, ?\Pi', \perp_2}{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), \Pi, ?\Pi', \perp_1, \perp_2} \\
\vdots \\
\frac{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2}{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2} \\
\vdots \\
\frac{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2, \perp}{\vdash ?\Gamma, ?\Delta, ?(P \otimes Q), ?\Pi, ?\Pi', \perp_1, \perp_2, \perp \& \perp}
\end{array}
\end{array}$$

Remark : All the (proof-theoretical) solutions that we know to the problem of extracting the computational content from the proofs of classical logic *choose* as the translation of the rule $\wedge_m$ with two positive active formulas that are in the stoup *one* of the two linear subproofs premisses of the last rule introducing $\perp \& \perp$.

The upshot of this work is precisely *not to choose* but rather to superimpose. This corresponds to the basic idea that a proof in LK is a superimposition of proofs of linear logic (or of LC).

B.4 Canonical form

B.4.1. DEFINITION. We say that the rule R of LK_{pol} is difficult iff $R = \wedge_m$ and both the active formulas in the premisses of R are positive and are not in the stoup.

Remark : It is proven in [Girard 91b] and in [DJS 97] that each rule R of LK_{pol} can be translated by “some” rules of LC :

- if R is already an LC -rule then the translation is straightforward
- if R is not a rule of LC (and if it is not difficult), then one cuts the premise(s) of the rule with a piece of η -proof of the conclusion of R . This procedure (which is also called a way of “constraining” R) might involve a choice in the order of performance of two cut rules ; but this choice has no computational content : the two proofs corresponding to the two possible choices are interpreted by the same clique in the P -semantics

(iii) if R is a difficult rule, then one translates it like in case (ii), but now the order of performance of the two cuts *does* matter : not only the two proofs corresponding to the two possible choices are interpreted by different cliques in the P -semantics, but it might very well be the case that they cannot be identified by *any* (non degenerated) denotational semantics.

One has two different ways of constraining each difficult rule, which correspond exactly to the two different ways of translating R in LL through the P -embedding.

This means (as expected) that there exists a classical version of P^a (which will be called C^a) that translates proofs of LK_{pol} into proofs of LC . It is the identity on formulas, and acts on sequents and proofs like P^a : just substitute “ $\neg V$ ” for “ $\perp$ ”, “ $\wedge_a$ ” for “ $\&$ ”, and erase the exponential rules (through C^a we are constraining the proofs of LK_{pol} in “all the possible different ways”). For any LK_{pol} -proof α , $P^a(\alpha)$ and $P(C^a(\alpha))$ have the same denotational semantics.

We decided to focus on P^a (rather than on C^a) simply to stress the relation with LL .

Remark : Let α be a proof of LK_{pol} with k difficult rules. Let's call α^c the set of LC -proofs obtained from α by constraining all its rules in all the possible *different* ways. From the previous remark we deduce that there are 2^k such LC -proofs (because there are two different ways of constraining each difficult rule of α).

The reader should note that α^c is uniquely determined only from the *semantical* point of view : to each LK_{pol} -proof α is associated a unique set of 2^k cliques of the coherent space interpreting the conclusion of α . However, as stated in (ii) of the previous remark, there might be different syntactical representations of “the same” proof of α^c . From now on, α^c will denote any such set of LC -proofs.

B.4.2. DEFINITION. Let α be an LK_{pol} -proof with k difficult rules and $\alpha^c = \{\alpha_1, \dots, \alpha_{2^k}\}$. We'll call slice of α any LC -proof belonging to α^c .

B.4.3. DEFINITION. Let α be an LK_{pol} -proof with k difficult rules and $\alpha^c = \{\alpha_1, \dots, \alpha_{2^k}\}$.

For all $i \in \{1, \dots, 2^k\}$, we define $P^+(\alpha_i)$ as the proof $P(\alpha_i)$, to which we added k occurrences of the rule introducing $\perp$.

B.4.4. DEFINITION. Let α be an LK_{pol} -proof with conclusion $\vdash \Gamma; \Pi$ and k difficult rules $R_1, \dots, R_k$, and let $\alpha^c = \{\alpha_1, \dots, \alpha_{2^k}\}$.

To each $i \in \{1, \dots, 2^k\}$ we can associate a sequence $\perp_{i_1}^1, \dots, \perp_{i_k}^k$, where for all $j \in \{1, \dots, k\}$ $i_j \in \{1, 2\}$, in such a way that R_j is constrained in the same way in α_l and α_m iff $l_j = m_j$.

This gives a bijection between $\{P^+(\alpha_i) : i \in \{1, \dots, 2^k\}\}$ and $\{\perp_{i_1}^1, \dots, \perp_{i_k}^k : i_j \in \{1, 2\}\}$.

We define the proof $\&(P^+(\alpha_i))$ of LL by building successively :

- the 2^{k-1} proofs with conclusion $\vdash ?\Gamma, \Pi, \perp_1^1 \& \perp_2^1, \perp_{i_2}^2, \dots, \perp_{i_k}^k$, where $1 \leq i_j \leq 2$
- the 2^{k-2} proofs with conclusion $\vdash ?\Gamma, \Pi, \perp_1^1 \& \perp_2^1, \perp_1^2 \& \perp_2^2, \perp_{i_3}^3, \dots, \perp_{i_k}^k$, where $1 \leq i_j \leq 2$

⋮

- the proof $\&(P^+(\alpha_i))$ with conclusion $\vdash ?\Gamma, \Pi, \perp_1^1 \& \perp_2^1, \perp_1^2 \& \perp_2^2, \dots, \perp_1^{k-1} \& \perp_2^{k-1}, \perp_1^k \& \perp_2^k$.
(The following example, in case $k = 2$, should clarify the construction).

Example : Suppose that in the previous definition $k = 2$. Then we have $\alpha^c = \{\alpha_1, \alpha_2, \alpha_3, \alpha_4\}$, and for $i \in \{1, 2, 3, 4\}$ we have that $P^+(\alpha_i)$ is an LL-proof with conclusion $\vdash ?\Gamma, \Pi, \perp_{i_1}^1, \perp_{i_2}^2$ where $i_1, i_2 \in \{1, 2\}$. Suppose for example that the conclusion of $P^+(\alpha_1)$ is $\vdash ?\Gamma, \Pi, \perp_1^1, \perp_1^2$, the conclusion of $P^+(\alpha_2)$ is $\vdash ?\Gamma, \Pi, \perp_1^1, \perp_2^2$, the conclusion of $P^+(\alpha_3)$ is $\vdash ?\Gamma, \Pi, \perp_2^1, \perp_1^2$, the conclusion of $P^+(\alpha_4)$ is $\vdash ?\Gamma, \Pi, \perp_2^1, \perp_2^2$.

Then, to obtain the proof $\&(P^+(\alpha_i))$, we build first the 2 ($= 2^{2-1}$) proofs $P^+(\alpha_1) \& P^+(\alpha_3)$ with conclusion $\vdash ?\Gamma, \Pi, \perp_1^1 \& \perp_2^1, \perp_1^2$ and $P^+(\alpha_2) \& P^+(\alpha_4)$ with conclusion $\vdash ?\Gamma, \Pi, \perp_1^1 \& \perp_2^1, \perp_2^2$, and then the proof $\&(P^+(\alpha_i))$ with conclusion $\vdash ?\Gamma, \Pi, \perp_1^1 \& \perp_2^1, \perp_1^2 \& \perp_2^2$.

Remark : In the previous definition, it is crucial to observe that, in the construction of the proof $\&(P^+(\alpha_i))$, the formulas active in the additive contractions of the rules $\&$ (that is the context formulas) which are different from $?\Gamma, \Pi$ are $\perp_1^i \& \perp_2^i$, or $\perp_1^i$, or $\perp_2^i$.

The fact that there are no additive contractions between $\perp_l^i$ and $\perp_m^i$ with $l \neq m$ means that there is no “mixing” between two “ $\perp$ ” which stem from two different ways of constraining the rule R_i .

Observe also that all the linear proofs $\&(P^+(\alpha_i))$ associated to the LK_{pol} -proof α are semantically equal.

B.4.5. PROPOSITION. (canonical form)

Let α be an LK_{pol} -proof, let $R_1 \dots R_k$ be its difficult rules, and let $\alpha^c = \{\alpha_1, \dots, \alpha_{2^k}\}$.

Then $[P^a(\alpha)] = [\&(P^+(\alpha_i))]$, where if d is a linear proof then $[d]$ is its denotational semantics.

proof : The idea is that, because the formulas $\perp \& \perp$ introduced by the embedding P^a are all formulas of the sequent conclusion of $P^a(\alpha)$ and because they are reversible and purely negative, it is possible to apply the rules introducing them “at the end” of the proof. The main point is of course the commutation of the rules introducing $\perp \& \perp$ with the promotion rule. But remember that the rule $\perp$ commutes to the promotion, and that $\&$ is reversible in a very strong sense : the denotational semantics of the proof obtained by applying a structural operation to the formula $A \& B$ (where A and B are linear formulas of type “?”) is the denotational semantics of the proof obtained by applying first the same structural operation to A and to B and then by applying the rule introducing $\&$ (see [Girard 91b] and [QuatrTorto 96] for more details).

A detailed proof can be given by induction on the proof α , i.e. by considering all the possible cases for the last rule R of α . There is nothing special about this induction, so that we just give an example : suppose that R is a binary logical rule which is also a rule of LC , and let γ and δ be the two subproofs of α that are premisses of R . If there are h (resp. g) difficult rules in γ (resp. δ), then $k = h + g$. Let α^c be obtained from the two sets $\gamma^c = \{\gamma_1, \dots, \gamma_{2^h}\}$ and $\delta^c = \{\delta_1, \dots, \delta_{2^g}\}$ by applying the rule R to γ_i and δ_j where $1 \leq i \leq 2^h$ and $1 \leq j \leq 2^g$. The induction hypothesis gives : $[P^a(\gamma)] = [\&(P^+(\gamma_1), \dots, P^+(\gamma_{2^h}))]$ and $[P^a(\delta)] = [\&(P^+(\delta_1), \dots, P^+(\delta_{2^g}))]$. Let's call R^* the P^a -translation of R (note that R^* can include one or two promotion rules). We have that $[P^a(\alpha)] = [R^*\{\&(P^+(\gamma_1), \dots, P^+(\gamma_{2^h})), \&(P^+(\delta_1), \dots, P^+(\delta_{2^g}))\}]$, where $R^*\{\&(P^+(\gamma_1), \dots, P^+(\gamma_{2^h})), \&(P^+(\delta_1), \dots, P^+(\delta_{2^g}))\}$ means that we apply R^* to the proofs

$\&(P^+(\gamma_1), \dots, P^+(\gamma_{2^h}))$, and $\&(P^+(\delta_1), \dots, P^+(\delta_{2^g}))$. After reversion of the rules introducing $\perp \& \perp$, we obtain $[P^a(\alpha)] = [\&(P^+(\alpha_1), \dots, P^+(\alpha_{2^k}))]$. $\square$

B.4.6. DEFINITION. Let α be a proof of LK_{pol} with k difficult rules and let $\alpha^c = \{\alpha_1, \dots, \alpha_{2^k}\}$.

We shall call canonical form of α , and we shall denote by $\bigwedge_a(\alpha_1, \dots, \alpha_{2^k})$ the proof of LC obtained by applying to $\{\alpha_1, \dots, \alpha_{2^k}\}$ exactly the same construction applied in definitions B.4.3 and B.4.4 to $\{P(\alpha_1), \dots, P(\alpha_{2^k})\}$: just substitute $\neg V$ for $\perp$ and $\bigwedge_a$ for $\&$.

B.4.7. COROLLARY. Let α be a proof of LK_{pol} with k difficult rules and let $\alpha^c = \{\alpha_1, \dots, \alpha_{2^k}\}$.
Then $[P^a(\alpha)] = [P(\bigwedge_a(\alpha_1, \dots, \alpha_{2^k}))]$ (or, otherwise stated, $[P(C^a(\alpha))] = [P(\bigwedge_a(\alpha_1, \dots, \alpha_{2^k}))]$).

proof : It is the LC -counterpart of the canonical form theorem (proposition B.4.5). $\square$

Remark : The corollary says (in particular) that all the canonical forms of an LK_{pol} -proof α have the same denotational semantics.

B.5 Cut-elimination in LK_{pol}

B.5.1. DEFINITION. Let α and β be two proofs in LK_{pol} . We will say that β is obtained from α by applying one cut-elimination step, and we will write $\alpha \rightarrow \beta$, iff there exist a canonical form $\bigwedge_a(\alpha_1, \dots, \alpha_{2^k})$ of α and a canonical form $\bigwedge_a(\beta_1, \dots, \beta_{2^h})$ of β s.t. :

- $k = h$
 - there exists $i \in \{1, \dots, 2^k\}$ such that in LC $\alpha_i \rightarrow \beta_i$ (that is β_i is obtained from α_i by applying one cut-elimination step of LC)
 - $\forall j \in \{1, \dots, 2^k\}, j \neq i$, we have $\alpha_j = \beta_j$.
- $\rightarrow$ will denote, as usual, the reflexive and transitive closure of $\rightarrow$.

An LK_{pol} -proof is said to be normal when it is an LC -proof without cuts.

Remark : (i) All the reducts of a proof α of LK_{pol} can be “materialized” only through their canonical forms.

(ii) If α is an LK_{pol} -proof without cuts, then its canonical forms are not necessary cut-free (take any proof without cuts containing a rule which is not an LC -rule). Intuitively, this means that such a proof is not (yet) completely explicit.

B.5.2. THEOREM. If in LK_{pol} $\alpha \rightarrow \beta$, then $[P^a(\alpha)] = [P^a(\beta)]$, i.e. P^a is a denotational semantics for LK_{pol} .

proof : To prove that P^a induces a denotational semantics for LK_{pol} , it is clearly enough to prove that if

$\alpha \rightarrow \beta$, then $[P^a(\alpha)] = [P^a(\beta)]$. But remember that it is proven in [QuatrTorto 96] that P is a denotational semantics for LC , so that from corollary B.4.7 we deduce : $[P^a(\alpha)] = [P(\bigwedge_a(\alpha_1, \dots, \alpha_{2^k}))] = [P(\bigwedge_a(\beta_1, \dots, \beta_{2^k}))] = [P^a(\beta)]$.

$\square$

Remark : Observe that, because to an LK_{pol} -proof α we can associate (syntactically) different sets α^c , the cut-elimination procedure previously defined does not enjoy the Church-Rosser property. For the same reason, although cut-elimination in LK_{pol} seems to be strongly normalizing, this does not follow immediately from the same result for LC proven in [DJS 97] and [JST 95].

This seems to indicate that (following a suggestion of [Girard 91b]), one should perhaps try to find a better syntax (which would identify two different ways of representing “the same” proof of α^c) for classical logic.

B.5.3 Gentzen ’s cut-elimination is not explication

This is an informal subsection, where we give two examples, the aim of which is simply to suggest a possible retrospective analysis of Gentzen ’s procedure of cut-elimination in the framework of LK_{pol} . What turns out is that this procedure is not (in general) compatible with the P^a -denotational semantics. Intuitively, if we assume that $P^a(\alpha)$ is the implicit content of the LK_{pol} -proof α , then Gentzen’s cut-elimination does something different from simply making explicit what is already implicit in the proof. Sometimes it interacts with explication by making choices (like in example 1 : it “does too much”), and sometimes it adds some “new” implicit content (like in example 2 : “it doesn’t do enough”) : in both cases this procedure deeply modify the computational content of the proof.

example 1 : The proof α (where R is supposed to be the unique difficult rule)

$$\frac{\frac{\frac{\beta}{\vdots} \quad \frac{\gamma}{\vdots}}{\vdash P; \Pi} \quad \frac{\vdash Q; \Pi'}{\vdash P \wedge_m Q, \Pi, \Pi';} \quad R \quad \frac{\frac{\delta}{\vdots}}{\vdash \neg P, \neg Q; \Pi''} \quad \frac{\vdash \neg P \vee_m \neg Q; \Pi''}{\vdash \Pi, \Pi', \Pi'';}$$

is substituted by one of the two following proofs α_1 and α_2 :

$$\begin{array}{c} \frac{\frac{\beta}{\vdots} \quad \frac{\gamma}{\vdots}}{\vdash P; \Pi} \quad \frac{\frac{\delta}{\vdots}}{\vdash \neg P, \neg Q; \Pi''} \\ \hline \frac{\vdash Q; \Pi' \quad \vdash \neg P, \neg Q; \Pi''}{\vdash \neg P, \Pi'', \Pi'} \\ \hline \frac{\vdash \Pi', \Pi'', \Pi}{\vdash \Pi, \Pi', \Pi'';} \end{array} \quad \begin{array}{c} \frac{\gamma}{\vdots} \quad \frac{\beta}{\vdots} \quad \frac{\delta}{\vdots} \\ \hline \frac{\vdash Q; \Pi' \quad \vdash \neg P, \neg Q; \Pi''}{\vdash \neg Q, \Pi'', \Pi} \\ \hline \frac{\vdash \Pi, \Pi'', \Pi'}{\vdash \Pi, \Pi', \Pi'';} \end{array}$$

It is possible to check that each of these two proofs is a reduct (w.r.t. the cut-elimination of LC) of one of the two slices of α ; choosing one of them means interacting with the process of cut-elimination of LK_{pol} . In LL this interaction is a cut between $\&(P^+(\alpha_1), P^+(\alpha_2))$ and one of the (two) proofs of $1 \oplus 1$ (which “selects” one of the two slices of the reduct of α).

example 2 : The proof α

$$\frac{\frac{\frac{\vdash \neg R; R \quad \vdash \neg R; R}{\vdash R, R; \neg R \wedge_m \neg R}}{\vdash R; \neg R \wedge_m \neg R} \quad \frac{\beta}{\vdash \neg R, P \wedge_m Q;}}{\vdash P \wedge_m Q; \neg R \wedge_m \neg R}$$

is substituted by the proof α'

$$\frac{\frac{\frac{\beta}{\vdash \neg R, P \wedge_m Q;}}{\vdash P \wedge_m Q, P \wedge_m Q; \neg R \wedge_m \neg R} \quad \frac{\frac{\beta}{\vdash \neg R, P \wedge_m Q;}}{\vdash P \wedge_m Q, P \wedge_m Q; \neg R \wedge_m \neg R}}{\vdash P \wedge_m Q; \neg R \wedge_m \neg R}$$

Suppose that $P \wedge_m Q$ in the subproof β is the conclusion of the unique difficult rule of α ; then α has two slices while α' has four (because it has *two* difficult rules).

Observe also that the two slices of α reduce to two (of the four) slices of α' , and that the slices of α' which are not reducts of the slices of α cannot (in general) be identified with the reducts of the two slices of α without producing (algorithmic) inconsistency : to identify them means to identify all the proofs of a given sequent.

The point is that for β we have two possible choices for the translation of the difficult rule (let 's call them $P(\beta_{PQ})$ and $P(\beta_{QP})$) and in $P^a(\alpha)$ when we reduce the cut following the contraction on $?P(R)$, we have to contract all the context formulas ... including $\perp \& \perp$ introduced by the translation! This means that the reduct of $P^a(\alpha)$ choose either $P(\beta_{PQ})$ or $P(\beta_{QP})$ for *both* the copies of β , while the P -image of one of the slices of α' will contain $P(\beta_{PQ})$ and $P(\beta_{QP})$ as subproofs.

Actually, for α' to have the same implicit content of α , one has to introduce the notion of “synchronization” of two difficult rules, and to synchronize the two difficult rules of α' . Technically speaking, to synchronize two difficult rules (which of course don't need to be equal, i.e. to have the same active formulas and conclusions) means precisely to contract the two “ $\perp \& \perp$ ” introduced by the P^a -translation of the two rules.

We see that in this case Gentzen's procedure “doesn't do enough” to keep the same implicit content.

B.6 Computational isomorphisms in LK_{pol}

B.6.1. DEFINITION. We say that R and Q (resp. M and N) are (ϕ, ψ) -isomorphic in LK_{pol} , and we write $R \stackrel{\phi, \psi}{\simeq} Q$, (resp. $M \stackrel{\phi, \psi}{\simeq} N$) if there exist a proof ϕ without difficult rules of $\vdash \neg R; Q$ (resp. a proof ϕ without difficult rules of $\vdash M; \neg N$) and a proof ψ without difficult rules of $\vdash \neg Q; R$ (resp. a proof ψ without difficult rules of $\vdash N; \neg M$) s. t.

- the coherent semantics of the proof $cut(\phi, \psi)$ of $\vdash \neg R; R$ (resp. $\vdash M; \neg M$) is the identity of $P(R)$ (resp. of $P(M)$)
- the coherent semantics of the proof $cut(\psi, \phi)$ of $\vdash \neg Q; Q$ (resp. $\vdash N; \neg N$) is the identity of $P(Q)$ (resp. of $P(N)$).

Remark : We ask positive formulas in the conclusions of ϕ and ψ to be in the stoup, so that the cuts between ϕ and ψ are interpreted as composition of functions in the coherent semantics. This ensures us that when $A \stackrel{\phi, \psi}{\simeq} B$, then $P(A)$ and $P(B)$ are isomorphic as coherent spaces.

B.6.2. THEOREM. In LK_{pol} we have the same computational isomorphisms than in LC .

proof : Girard proves (in [Girard 91b]) that the correlation semantics of LC satisfies a certain number of isomorphisms.

We have already proven (in chapter 15) that these isomorphisms are also satisfied if one takes the (slightly more syntactical) previous definition of isomorphism, and simply by considering the usual coherent semantics induced by the P -embedding (see [QuatrTorto 96]).

The fact that this still holds for LK_{pol} (which contains LC as a subsystem) is straightforward : the definition of isomorphism is the same! $\square$

Remark : (i) Observe however that, in spite of the previous theorem, the situation in LK_{pol} is different from the one we have in LC . Consider for example the associativity of conjunction.

If we take three proofs α_1 of $\vdash P$; α_2 of $\vdash Q$; and α_3 of $\vdash R$; then the two proofs β of $\vdash (P \wedge_m Q) \wedge_m R$; and γ of $\vdash P \wedge_m (Q \wedge_m R)$; that one can get simply by performing the two conjunctions in a different order, have both four slices. Let's distinguish them by the order of application of the

cut-rules of LK_{pol} between the pieces of η -proofs and the proofs α_1 , α_2 and α_3 .

The slice QPR (for example) of β is not a slice of γ . But this is only apparently contradictory with the isomorphism $(P \wedge_m Q) \wedge_m R \stackrel{\phi, \psi}{\simeq} P \wedge_m (Q \wedge_m R)$. Actually, roughly speaking, there exists a proof of

$\vdash P \wedge_m (Q \wedge_m R)$; containing “the same” slices of β . To convince himself, the reader may observe that what one eventually does to get the proof β of $\vdash (P \wedge_m Q) \wedge_m R$; from α_1 , α_2 and α_3 is to cut α_1 , α_2 and α_3 with the canonical proof of $\vdash \neg P, \neg Q, \neg R; (P \wedge_m Q) \wedge_m R$ in four of the six possible ways. Then, one obtains a proof β' of $\vdash P \wedge_m (Q \wedge_m R)$; with “the same” slices of β by cutting α_1 , α_2 and α_3 with the canonical proof of $\vdash \neg P, \neg Q, \neg R; P \wedge_m (Q \wedge_m R)$ in *the same* four ways (which *are not* the four ways used to obtain the four slices of γ).

This last proof β' is precisely a reduct of the proof obtained by cutting β with the (canonical) proof ϕ of $\vdash \neg(P \wedge_m Q) \vee_m \neg R; P \wedge_m (Q \wedge_m R)$.

(ii) If one really wants to have simultaneously, for the proofs of $\vdash P \wedge_m (Q \wedge_m R)$; and of $\vdash (P \wedge_m Q) \wedge_m R$; obtained from α_1 , α_2 and α_3 of (i), all the six possible slices, then one has to take a version of LK_{pol} with difficult rules of arbitrary arity.

But observe that this is not necessary to have a computational isomorphism between $A \wedge_m (B \wedge_m C)$ and $(A \wedge_m B) \wedge_m C$, whatever the polarities of A , B and C are.

B.7 Further developpements

From the technical point of view, the next step (which looks difficult) would be to find a better syntax for LK_{pol} , identifying all the canonical forms of a given proof.

What seems interesting (in the framework of LK_{pol}) is to define some kind of relationship between the slices superimposed by an LK_{pol} -proof.

The notion of synchronization of two (or more) difficult rules might also play an important role. It is a very natural notion, which is technically translated by one (or more) contraction(s) on the $\perp \& \perp$ introduced by P^a (as stated in example 2 of 5.1).

To take difficult rules of arbitrary arity might also increase the flexibility of the calculus.

Finally, following an old idea of [Girard 91b], one might also try to “understand” Gentzen’s cross-cut procedure in the framework of LK_{pol} .