



Géométrie de l'Interaction et Réseaux Différentiels

Marc de Falco

► To cite this version:

Marc de Falco. Géométrie de l'Interaction et Réseaux Différentiels. Mathématiques [math]. Université de la Méditerranée - Aix-Marseille II, 2009. Français. NNT: . tel-00392242

HAL Id: tel-00392242

<https://theses.hal.science/tel-00392242>

Submitted on 6 Jun 2009

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

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

UNIVERSITÉ DE LA MÉDITERRANÉE, AIX-MARSEILLE 2
U.F.R. DE MATHÉMATIQUES

THÈSE

pour obtenir le grade de
DOCTEUR DE L'UNIVERSITÉ AIX-MARSEILLE 2

Discipline : Mathématiques

présentée et soutenue publiquement par

Marc de Falco

le 28 mai 2009

Titre :

Géométrie de l'Interaction et Réseaux Différentiels

Directeur de thèse : M. Laurent REGNIER

Rapporteurs : M. Jean GOUBAULT-LARRECQ
M. Stefano GUERRINI

Jury : M. Yves LAFONT, *Président*
M. Thomas EHRHARD
M. Georges GONTHIER
M. Jean GOUBAULT-LARRECQ
M. Stefano GUERRINI
M. Laurent REGNIER
M. Philip SCOTT

Table des matières

Remerciements	5
Introduction	7
I Réseaux d'interaction et logique	15
1 Préliminaires	17
1.1 Présentation traditionnelle des réseaux d'interaction	17
1.2 Permutations et injections partielles	19
2 Réseaux d'interaction explicites	25
2.1 Motivations	25
2.2 Statique	26
2.3 Opérations sur les réseaux d'interaction	29
2.4 Dynamique	32
2.5 Morphologie générale des réseaux	33
2.6 Les réseaux d'interaction sont le Ex-effondrement des réseaux Axiome/Coupure	33
2.7 La catégorie $\mathbf{IN}$ et l'approche DPO	36
3 Boîtes	43
3.1 Définition	43
3.2 Opérations de base sur les boîtes	44
3.3 Boîtes additives	47
4 Réseaux d'interaction et logique linéaire	49
4.1 Bibliothèques closes	49
4.2 Typage de réseaux	49
4.3 MLL	50
4.4 MELL	51
II Géométrie de l'Interaction	55
5 Semigroupes inversifs	57
5.1 Définitions	57
5.2 Propriétés essentielles	61

6	Chemins et réduction	67
6.1	Graphes de réseaux	67
6.2	Chemins	67
6.3	Réduction de chemins	69
7	Géométrie de l'Interaction	75
7.1	Persistance et pondérations fidèles	75
7.2	Géométrie de l'Interaction	79
7.3	Exemples de GdI	80
8	λ-calcul	83
8.1	Définitions	84
8.2	Exemples	84
8.3	La KAM	85
8.4	Traduction du λ -calcul dans SMELL	86
9	Concision de la GdI des λ-termes	91
9.1	Concision	91
9.2	Un calcul de $ \delta\overline{n} $	91
9.3	Conjectures	101
III	Les Réseaux d'Interaction Différentiels et leur Géométrie de l'Interaction	103
10	Réseaux d'interaction différentiels	105
10.1	Définitions	105
10.2	Interprétations des règles de réduction	109
10.3	Remarques sur la nature sémantique de la somme	113
11	Géométrie de l'interaction des réseaux d'interaction différentiels	115
11.1	Chemins dans des sommes	115
11.2	L'algèbre ∂L^*	123
11.3	Pondération des réseaux d'interaction différentiels	125
11.4	Une réalisation de $\partial L^* : \mathfrak{dL}$	132
11.5	Graphes de superposition	134
	Perspectives	139
	Index	140
	Bibliographie	143

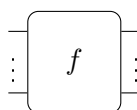
Remerciements

Introduction

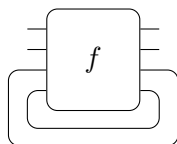
Le programme de *Géométrie de l'Interaction*

La rétroaction et la programmation

Le phénomène de rétroaction, en anglais *feedback*, est présent dans quasiment toutes les disciplines scientifiques qui décrivent l'évolution de systèmes : *physique, biologie, automatique, informatique, ...* Schématiquement, il s'agit de considérer un processus de transformation vu comme une boîte noire possédant un certain nombre d'entrées et de sorties :



et de brancher des sorties sur des entrées, créant ainsi des boucles dites de *rétroaction* :



On définit ainsi un nouveau processus de transformation possédant moins d'entrées et moins de sorties. Du point de vue de l'étude de ces processus, le nombre de degrés de liberté a nécessairement diminué. La rétroaction effectue donc une réduction par *repli* d'un système sur lui-même.

Dans le cadre de l'informatique, les processus de transformation que l'on manipule sont des programmes. L'exécution de ces programmes est alors un processus externe permettant d'aboutir à une expression simple des sorties en fonction des entrées. Les inconvénients liés à cette externalisation sont clairs : tout est lié à un élément extérieur, la syntaxe des programmes n'a pas de lien interne avec ses propriétés dynamiques. C'est en partant de ce constat que Jean-Yves Girard énonce dans l'article [Gir89b] son programme de *Géométrie de l'Interaction* visant à définir l'exécution comme un processus de rétroaction. Bien entendu ce programme fut exprimable grâce à l'introduction de la *logique linéaire* et de la notion de réseaux de preuves. La réussite d'un tel programme a pour conséquence d'introduire une notion de syntaxe justifiée par le fait qu'elle contient sa propre réduction.

Logique linéaire et réseaux de preuves

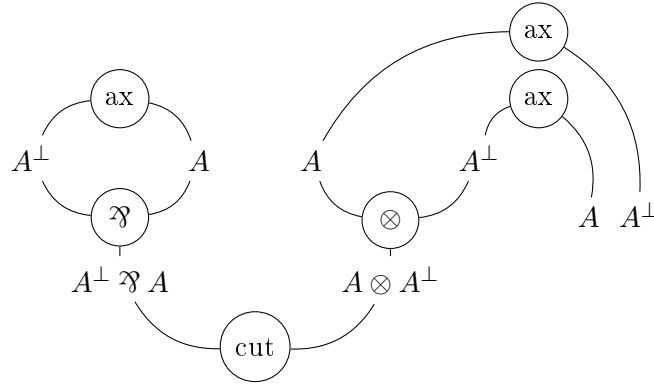
La logique linéaire provient d'un relèvement syntaxique d'une décomposition sémantique de l'implication intuitionniste. Dans le modèle des espaces cohérents l'implication intuitionniste

$A \rightarrow B$ se décompose en $!A \multimap B$ où $!A$ correspond à la répétition de A et $\multimap$ est une sorte de fonction linéaire. Dans la syntaxe cela se traduit par une dichotomie des formules : d'un coté des formules que l'on ne peut *utiliser* qu'une seule fois et de l'autre certaines que l'on peut utiliser un nombre quelconque de fois. La modalité $!$ permet de passer des premières aux secondes.

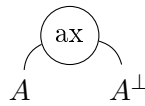
Les objets que manipule la logique linéaire ne sont donc plus les formules mais les occurrences. On peut ainsi associer à une preuve une sorte de graphe sur les occurrences la représentant. Ainsi à la preuve

$$\frac{\frac{\frac{\overline{\vdash A^\perp, A}}{\vdash A^\perp \wp A} \text{ax}}{\vdash A^\perp \wp A} \wp \quad \frac{\frac{\frac{\overline{\vdash A, A^\perp}}{\vdash A \otimes A^\perp, A, A^\perp} \text{ax}}{\vdash A \otimes A^\perp, A, A^\perp} \otimes}{\vdash A, A^\perp} \text{cut}$$

on associera le graphe R



que l'on appelle réseau de preuve. Celui-ci peut être réduit par *élimination des coupures* et donne le réseau normal R_0



qui est le résultat final du programme représenté par la preuve précédente à travers la correspondance de Curry-Howard.

Dans ce contexte, l'objectif de la Géométrie de l'Interaction (*GdI*) est d'associer au réseau R précédent un objet mathématique

$$\begin{array}{c} A \text{ --- } \boxed{R^\bullet} \text{ --- } A^\perp \\ A^\perp \wp A \text{ --- } \boxed{R^\bullet} \text{ --- } A \otimes A^\perp \end{array}$$

tel que

$$\begin{array}{c} A \text{ --- } \boxed{R^\bullet} \text{ --- } A^\perp \\ \text{--- } \boxed{R^\bullet} \text{ --- } \end{array} = A \text{ --- } \boxed{R_0^\bullet} \text{ --- } A^\perp$$

où l'on a effectué un *calcul* de la rétroaction.

Mises en œuvre de la Géométrie de l'Interaction

L'approche algèbre d'opérateurs

L'approche initiale de Girard a été développée au cours d'une série d'articles. Tout d'abord, dans l'article [Gir87] il a observé que les preuves de la logique linéaire multiplicative pouvait être interprétée par deux permutations sur les conclusions : une permutation induite par les axiomes et une induite par les coupures. En se plaçant dans l'espace vectoriel libre sur les conclusions on peut associer des applications linéaires à ces permutations, et on peut alors pour passer à la forme normale utiliser la *formule d'exécution* :

$$\text{Ex}(\Pi, \sigma) = (1 - \sigma^2)(\Pi + \Pi\sigma\Pi + \dots + (\Pi\sigma)^n\Pi + \dots)(1 - \sigma^2)$$

où Π est l'application associée aux axiomes et σ aux coupures. Dans l'article [Gir89a] l'auteur étend cette approche au système $\mathbf{F}$, les preuves sont interprétées dans un espace de Hilbert séparable H et il y est démontré que, d'une part la formule précédente ne contient qu'un nombre fini de termes, c'est-à-dire un résultat de nilpotence, et d'autre part elle encode effectivement la réduction pour certaines preuves. L'interprétation des preuves est alors défini en terme d'opérateurs bornés et donc d'éléments d'une algèbre stellaire.

Le cas du λ -calcul pur est traité dans l'article [Gir90] par l'intermédiaire de la nilpotence faible.

L'extension à la logique linéaire toute entière, et plus particulièrement aux additifs, est donnée par l'article [Gir95]. Les opérateurs agissent maintenant sur un produit $H \otimes H$ et sont considérés modulo isomorphisme sur la seconde composante, introduisant ainsi les notions de *dialectes* et de *variants*.

Dans l'article [Gir06] Jean-Yves Girard pose et résout en terme d'opérateurs *l'équation de rétroaction*, donnant ainsi l'étendue mathématique possible des interprétations de preuve. Cette ouverture est poursuivie dans l'article [Gir08] où le cadre est maintenant un facteur d'une algèbre de von Neumann. Ces deux derniers articles se démarquent des précédents dans la mesure où les mathématiques qui y sont introduites deviennent le cadre potentiel d'un renouveau de la syntaxe, cet objectif n'étant pas encore achevé à l'heure actuelle.

Les graphes de partage et la réduction optimale

Parallèlement à l'introduction de la Géométrie de l'Interaction, Lamping a présenté dans l'article [Lam90] un algorithme pour calculer la réduction optimale de Lévy des λ -termes, c'est-à-dire la réduction effectuant le plus de partage de redex. Son algorithme y est présenté en terme de réécriture de graphes, appelés *graphes de partage*. La validation de son algorithme repose sur une sémantique des chemins proche des opérateurs définis dans la GdI appelée *sémantique de contexte*.

Le lien entre les graphes de partage et la GdI est complètement explicité dans l'article [AGL92a] par Abadi, Gonthier et Lévy. Les auteurs utilisent dans l'article [AGL92b] les techniques de Lamping pour définir une réduction des réseaux de preuve à l'aide de graphes de partage et ainsi sans les boîtes.

De nombreux travaux ont été alors développés autour des graphes de partage et de la sémantique de contexte. Citons le livre [AG98] qui présente de manière très complète cette approche.

Catégories monoïdales à trace

La sémantique donnée par la GdI n'étant pas invariante par réduction, elle ne rentre pas dans le cadre traditionnel de la sémantique dénotationnelle reposant sur la théorie des catégories. L'article [AJ92] d'Abramsky et Jagadeesan a tout d'abord rapproché la GdI de la théorie des domaines, présentant ainsi des liens concrets avec les graphes de flots et la rétroaction. Ce n'est qu'après l'article [JSV96] de Joyal, Street et Verity que la construction catégorique sous-jacente à l'article précédent pu être complètement explicitée. Cette construction repose sur une généralisation de la notion de trace partielle présente en mécanique quantique. Celle-ci consiste à déduire l'état d'un système A à partir d'un sur-état d'un système $A + B$ par l'intermédiaire d'un calcul de trace sur le *bloc* correspondant à B . Du point de vue catégorique cela correspond alors à avoir une transformation dinaturelle $Tr_{B, B^\perp}^{A, A^\perp} : C(A \otimes B, A^\perp \otimes B^\perp) \mapsto C(A, A^\perp)$.

Cela a permis à Abramsky d'effectuer une relecture de son approche et d'introduire dans l'article [Abr97] une notion de *situation de GdI*, cadre catégorique basé sur la trace pour interpréter la GdI de la logique linéaire multiplicative exponentielle.

Une justification complète de ce cadre a été donnée par Scott et Haghverdi dans l'article [HS06] où les définitions et résultats de l'article fondateur [Gir89a] ont entièrement été établis dans un cadre catégorique, ouvrant ainsi la voie à des modèles dépassant le cadre de l'algèbre d'opérateurs.

Pondération de chemins dans les réseaux

L'approche initiale en terme d'algèbre d'opérateurs associe un opérateur global à un réseau. En fait, cet opérateur se décompose en sous-opérateurs simples associés à des chemins dans le réseau. Cette approche a été initiée par Danos et Regnier dans leurs thèses respectives [Dan90, Reg92] et continuée à travers une série d'articles [DR93, DR95].

Les chemins *réguliers*, ceux pondérés par un opérateur non nul, sont alors le centre d'intérêt de la GdI. La validité qui s'exprimait alors par une préservation de la formule d'exécution se traduit maintenant par une équivalence entre la notion de régularité et la notion de persistance, c'est-à-dire de non destruction des chemins par réduction vers la forme normale. L'équivalence entre les chemins réguliers et les chemins issus des graphes de partage est explicitée dans l'article [ADLR94].

On peut transformer un réseau en un automate bi-déterministe tel que les exécutions possibles correspondent au chemins réguliers. C'est ce que l'on appelle également une *machine à jetons*. Une approche très bas niveau en est donné par Mackie ([Mac95]) pour un langage type PCF, où la machine est directement implémentable en termes d'instructions assembleurs. Dans l'article [Lau01] Laurent donne une machine à jetons pour la logique linéaire complète basée sur l'article [Gir95].

De la GdI comme analyse statique des programmes

La dernière approche citée a en fait des liens profonds avec la programmation. Dans un programme donné il apparaît une notion de chemins d'accès liant l'utilisation d'une variable au lieu qui l'a définie. Après réduction, lorsqu'une variable est remplacée par un terme, de nouveaux chemins d'accès sont définis. Il est très délicat de décrire cette transformation des chemins d'accès. La Géométrie de l'Interaction offre un cadre algébrique pour effectuer cette

description. Ainsi, à l'aide de la GdI il est possible de *calculer* ses chemins d'accès et d'extraire de l'information d'un programme sans l'exécuter, c'est-à-dire faire de l'analyse statique.

Réseaux d'interaction différentiels

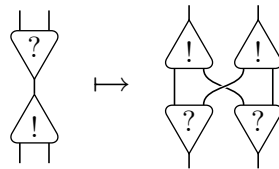
Dans une suite d'article Thomas Ehrhard a introduit des sémantiques quantitatives dans lesquelles apparaît naturellement une notion de différentiation des morphismes [Ehr02, Ehr05]. Procédant à un relèvement syntaxique de ces notions, à la manière de la linéarité dont le relèvement donna la logique linéaire, Ehrhard a introduit avec Regnier la notion de λ -calcul différentiel dans l'article [ER03].

Du point de vue logique la linéarité correspond à la notion de ressource utilisable une seule fois, donc un programme linéaire en un de ses arguments doit l'utiliser exactement une fois. Prenons maintenant un programme utilisant un nombre quelconque de fois son argument. Pour simplifier on considère un terme t de la forme $\lambda x.C[x, \dots, x]$ où $C[]$ est un contexte contenant un trou par occurrence de x liée au λ , et linéaire vis-à-vis de chacune de ces occurrences. La dérivée de t par rapport à x prise en u est la meilleure approximation linéaire de t prise en u . La seule manière de substituer linéairement u à x est de faire un choix parmi les occurrences de x , cela revient donc à prendre $C[x, \dots, u, \dots, x]$ où la i -ème occurrence de x a été substituée. Or, comme ce choix est arbitraire on introduit une notion syntaxique de somme de termes permettant de sommer sur i . On pose donc

$$(\lambda x.C[x_1, \dots, x_n])'(u) = \lambda x. \sum_i C[x_1, \dots, x_{i-1}, u, x_{i+1}, \dots, x_n]$$

Bien entendu le λx a été préservée car le terme résultant dépend toujours de x . En orientant cette égalité de gauche à droite on récupère un enrichissement de la β -réduction.

Cette approche peut être vu comme une déduction naturelle d'une extension de la logique linéaire : *la logique linéaire différentielle*. En effet, dans l'article [ER05] Ehrhard et Regnier ont introduit, sous forme de réseaux d'interaction, une telle extension. A la structure traditionnelle de co-monade issues des règles structurales, elle consiste à ajouter des règles co-structurales et, de ce fait, une structure de monade. Le tout forme désormais une structure de bigèbre. Le choix de la présentation en terme de réseaux d'interaction plutôt que de calcul des séquents ou de réseaux de preuve est due à cette structure de bigèbre. En effet, la règle de réduction



ne s'exprime pas aussi naturellement dans les autres formalismes.

La somme utilisée pour le calcul de la dérivée fait apparaître un possible lien avec le non-déterminisme. Ce lien est concrétisé dans l'article [EL07] où Ehrhard et Laurent présentent une traduction d'un fragment finitaire du π -calcul dans les réseaux d'interaction différentiels.

Organisation de la thèse

Le but de notre travail est d'étendre la Géométrie de l'Interaction aux réseaux d'interaction différentiels, apportant ainsi les prémisses d'une notion objective de réduction pour les calculs

de processus.

Dans un premier temps, il faut choisir la présentation adaptée de la GdI. L'approche algèbre d'opérateurs et l'approche catégorie à trace reposent sur une notion de type et une notion de preuve. Par exemple, les définitions et preuves de Girard dans ses articles s'effectuent à l'aide de raisonnements sur la dernière règle. Or, ici, la notion de preuve n'est pas directement accessible et le typage apparaît comme un frein à l'étude de possibles traductions exotiques dans les réseaux différentiels. L'approche graphes de partage semble difficile à mettre en place directement car il n'existe pas de notion de réduction optimale adaptée pouvant la guider.

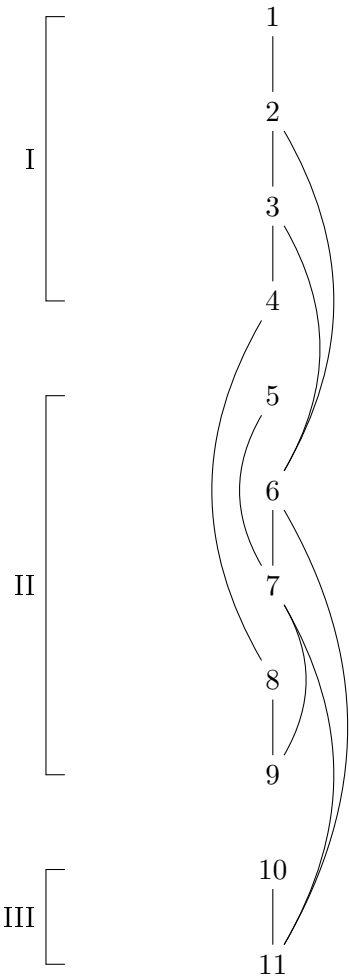
Reste donc l'approche pondération de chemins. Pour mettre en place celle-ci il est nécessaire d'introduire une notion de chemin adapté pour les réseaux d'interaction, d'en déduire une notion pour les réseaux d'interaction différentiels. Puis, il s'agit de trouver la pondération adéquate.

Pour pouvoir définir rigoureusement les chemins dans les réseaux d'interaction nous avons été amené à développer une variante complètement localisée de ceux-ci reposant sur la notion de permutations partielles. Elle est développée au chapitre 2, étendue aux boîtes au chapitre 3 et mise en relation avec les réseaux de preuves traditionnels au chapitre 4. La notion adaptée de chemin est alors étudiée au chapitre 6 et celle de pondération au chapitre 7.

Aux chapitres 8 et 9 nous revenons sur la GdI du λ -calcul en étudiant des termes pour lesquels elle n'est pas préservée par la réduction.

Ensuite, après avoir introduit les réseaux différentiels dans leur présentation usuelle au chapitre 10, nous présentons notre Géométrie de l'Interaction dans le chapitre 11.

Leitfaden



Première partie

Réseaux d'interaction et logique

Chapitre 1

Préliminaires

Nous allons présenter dans ce chapitre les définitions usuelles des réseaux d'interaction. Nous donnerons ensuite les ingrédients mathématiques principaux qui nous serviront à établir une présentation explicite des mêmes réseaux au chapitre suivant.

1.1 Présentation traditionnelle des réseaux d'interaction

Les définitions présentées reprennent pour l'essentiel celles de l'article originel d'Yves Lafont [Laf90].

Un *réseau d'interaction* est constitué d'un ensemble fini de ports libres, d'un ensemble fini de cellules contenant un symbole et une suite finie de ports, et d'une partition des ports en couples appelés fils. Les cellules sont représentées par des triangles, le premier de leurs ports, appelé principal, étant disposés sur l'un des sommets et les ports suivants, appelés auxiliaires, selon le côté opposé.



Un réseau peut également contenir le cas particulier d'un fil *branché sur lui-même* que l'on nomme boucle.



Un exemple de réseau d'interaction est présenté dans la figure 1.1.

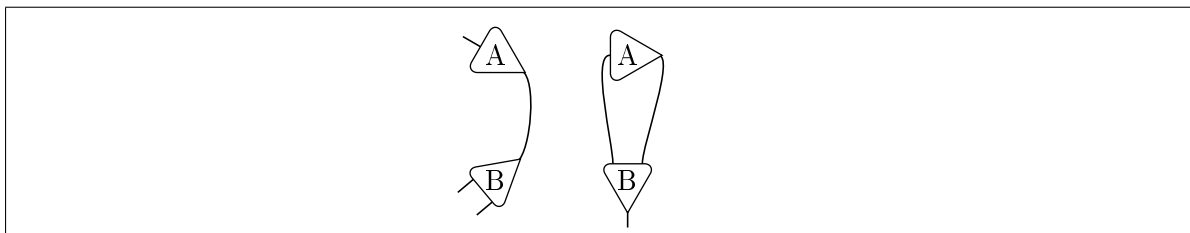
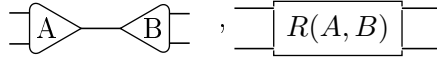


Figure 1.1: Un réseau d'interaction

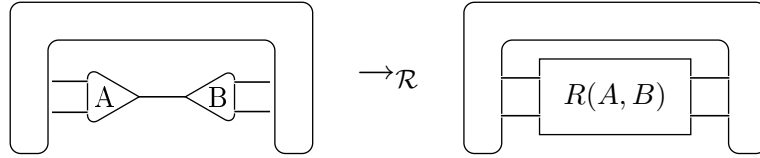
Les réseaux représentent la partie statique du formalisme, ils deviennent réellement intéressant dès lors qu'on leur adjoint une dynamique. Ils empruntent celle-ci à la réécriture de

graphe : à un sous-réseau particulier on va substituer un autre sous-réseau partageant la même *interface*. Une telle paire sera appelé une *règle*, son membre gauche un *redex* et son membre droit un *motif*. Sans autres précisions, cette définition paraphrase la réécriture de graphes. L'innovation des réseaux d'interactions est d'introduire une limitation sur les redex provenant des réseaux de preuve de la logique linéaire multiplicative. Les seuls redex autorisés correspondent à deux cellules connectées par leurs ports principaux, le fil les reliant étant appelé *fil actif*.

Ainsi, une règle $\mathcal{R}$ est un couple :

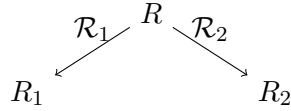


que l'on applique ainsi :

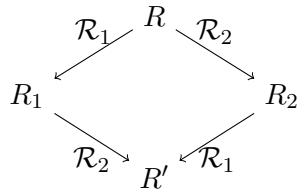


Une propriété essentielle des réseaux d'interactions est la confluence forte de la réduction.

Propriété 1.1 Soient R un réseau d'interaction, $\mathcal{R}_1$ et $\mathcal{R}_2$ des règles s'appliquant dans R sur des redex distincts telles que



Il existe un réseau R' tel que



Preuve La preuve de cette propriété découle directement du fait que deux redex distincts sont complètement déconnectés : ils sont définis par deux fils actifs différents. Étant donné que les réductions s'appliquent localement à l'endroit du redex, il est immédiat que l'on peut appliquer $\mathcal{R}_1$ et $\mathcal{R}_2$ indépendamment, ce qui produit le réseau R' voulu. $\blacktriangleleft$

Remarque 1.2 C'est ici que la simplification apportée par les réseaux d'interaction prend tout son sens. En effet, dans les systèmes de réécritures habituels, les redex peuvent se superposer, donnant lieu à des paires critiques, ou influencer les uns sur les autres, obligeant à raisonner sur des redex résiduels comme dans le λ -calcul.

On appelle un ensemble de règles, tel qu'il existe au plus une règle par paire de symboles, une *bibliothèque*. On note $\rightarrow_B$ la réduction associée à la bibliothèque B . La propriété 1.1 exprime la confluence forte de $\rightarrow_B$.

1.2 Permutations et injections partielles

Les définitions qui vont suivre sont usuelles dans le cadre des *petits modèles* de la *géométrie de l'interaction* [DR95] ou pour la définition de la catégorie monoïdale tracée **PInj** [HS06].

1.2.1 Permutations

On rappelle ici qu'on appelle permutation d'un ensemble E toute bijection de E dans lui-même, et l'on note $\mathfrak{S}_E$ leur ensemble. Étant donné $\sigma \in \mathfrak{S}_E$ on appelle *ordre de σ* le plus petit entier $n > 0$ tel que $\sigma^n = id_E$. Soit $x \in E$ on note $\text{Orb}_\sigma(x) = \{\sigma^i(x) \mid i \in \mathbb{N}\}$ et on l'appelle *l'orbite* de x . On note $\text{Orbs}(\sigma)$ l'ensemble des orbites de σ . Si o est une orbite on note $|o|$ sa taille.

On note $(c_1, \dots, c_n)$ la permutation envoyant c_i sur c_{i+1} , pour $i < n$, c_n sur c_1 et étant l'identité partout ailleurs, on l'appelle un *cycle de longueur n* , nombre qui est aussi son ordre. Toute permutation est un produit de cycles disjoints.

Soit σ une permutation de E et $\mathcal{L}$ un ensemble, on dit que σ est *étiquetée* par $\mathcal{L}$ si on a une fonction $l_\sigma : \text{Orbs}(\sigma) \rightarrow \mathcal{L}$. On dit que σ est *à orbites pointées* si elle est étiquetée par E et que $\forall o \in \text{Orbs}(\sigma)$ on a $l_\sigma(o) \in o$. On remarque qu'une orbite est un sous-cycle et ainsi, être à orbites pointées signifie que l'on a choisi un point de départ dans chacun de ces sous-cycles.

1.2.2 Injection partielles

Une *injection partielle (d'entiers)* f est une bijection d'un sous-ensemble de $\mathbb{N}$, appelé son *domaine* et noté $\text{dom}(f)$, sur un sous-ensemble de $\mathbb{N}$, appelé son *co-domaine* et noté $\text{codom}(f)$. On note $f : A \twoheadrightarrow B$ pour dire que f est une injection partielle de domaine A et de co-domaine B . On note f^* l'inverse de f vu comme une injection partielle de B dans A . On appelle *permutation partielle* une injection partielle f telle que $\text{dom}(f) = \text{codom}(f)$.

1.2.3 Exécution

Soit f une injection partielle et $E', F' \subseteq \mathbb{N}$. On note $f \upharpoonright_{F'}^{E'}$ l'injection partielle de domaine $\{x \in E' \cap \text{dom} f \mid f(x) \in F'\}$ et telle que $f \upharpoonright_{F'}^{E'}(x) = f(x)$ lorsque cela est défini. On a

$$f \upharpoonright_{F'}^{E'} : f^{-1}(F') \cap E' \twoheadrightarrow f(E') \cap F'$$

Si $E' = F'$ on note $f \upharpoonright_{E'} = f \upharpoonright_{E'}^{E'}$.

Lorsque $\text{dom}(f) \cap \text{dom}(g) = \emptyset$ et $\text{codom}(f) \cap \text{codom}(g) = \emptyset$, on dit que f et g sont *disjointes* et on définit la somme $f + g$, de domaine la réunion des domaines et à valeurs dans la réunion des co-domaines. On a naturellement un ordre associé à cette notion de somme : $f \leq g \iff \exists h, f + h = g$.

Propriété 1.3 Soient $f : A \uplus B \twoheadrightarrow C \uplus D$ et $g : D \twoheadrightarrow B$, situation représentée par le

$$\text{diagramme} \quad A \uplus B \begin{array}{c} \xrightarrow{f} \\ \xleftarrow{g} \end{array} C \uplus D \quad .$$

i) Pour tout $n \in \mathbb{N}$, l'injection partielle de A dans C

$$\text{Ex}_n(f, g) = f \upharpoonright_C^A + (fgf) \upharpoonright_C^A + \dots + (f(gf)^n) \upharpoonright_C^A$$

est bien définie.

- ii) $(\text{Ex}_n(f, g))_{n \in \mathbb{N}}$ est une suite croissante d'injections partielles, dont la limite est notée $\text{Ex}(f, g)$, on l'appelle l'exécution de f selon g .
- iii) Si $\text{dom}(f)$ est fini la suite $(\text{Ex}_n(f, g))_n$ est stationnaire et

$$\text{Ex}(f, g) : A \rightarrow C$$

La figure 1.2 donne une représentation graphique de l'exécution.

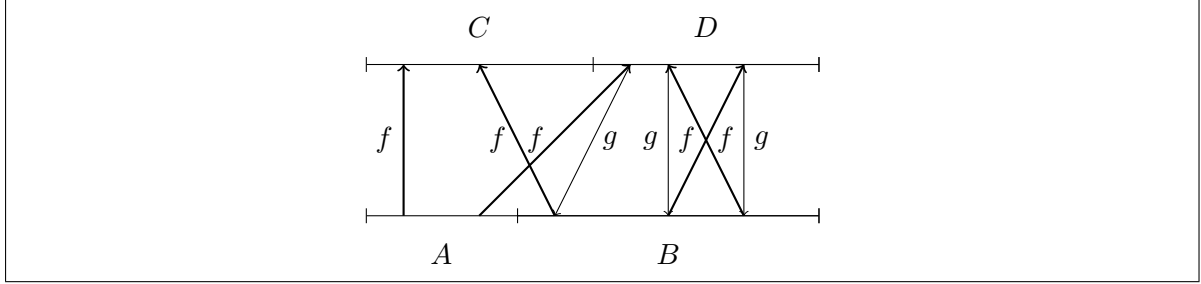


Figure 1.2: Représentation de $\text{Ex}(f, g)$ avec les notations de la propriété 1.3

Preuve i) Pour affirmer que la somme est valide il suffit de montrer que $\forall i \neq j \in \mathbb{N}$:

$$\begin{aligned} (f(gf)^i)(A) \cap (f(gf)^j)(A) \cap C &= \emptyset \\ (f(gf)^i)^{-1}(C) \cap (f(gf)^j)^{-1}(C) \cap A &= \emptyset \end{aligned}$$

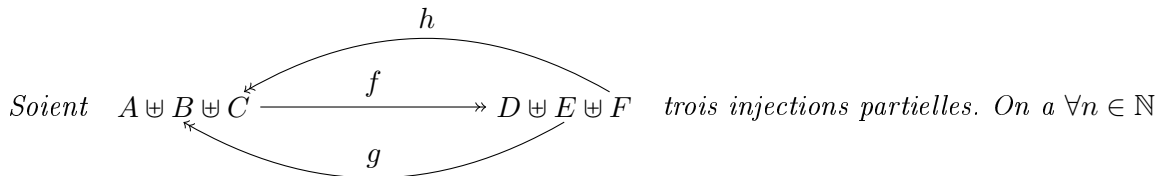
Supposons qu'il existe un $x \in (f(gf)^i)(A) \cap (f(gf)^j)(A) \cap C$, et soient y et $z \in A$ tels que $x = f(gf)^i(y) = f(gf)^j(z)$. On peut supposer de surcroît que $i < j$, et l'on en déduit alors que $y = (gf)^{j-i}(z) \in B$, ce qui est contradictoire dans la mesure où $y \in A$ et $A \cap B = \emptyset$.

L'autre égalité se montre de manière similaire.

ii) Soient $n \leq m \in \mathbb{N}$ et $x \in \text{dom}(\text{Ex}_n(f, g))$, par définition de la somme il existe un unique k tel que $\text{Ex}_n(f, g)(x) = (f(gf)^k)(x)$. Mais alors $x \in \text{dom}(\text{Ex}_m(f, g))$ et l'unicité de k entraîne que $\text{Ex}_m(f, g)(x) = (f(gf)^k)(x)$. Ainsi, $\text{Ex}_m(f, g)$ est une extension de $\text{Ex}_n(f, g)$.

iii) Supposons qu'il existe un $x \in A - \text{dom}(\text{Ex}(f, g))$, alors on devrait avoir, pour tout k , $(f(gf)^k)(x) \in D$ ou sinon $\text{Ex}(f, g)(x)$ serait défini. Mais D étant fini, il existe nécessairement deux entiers $n \leq m$ tels que $(f(gf)^n)(x) = (f(gf)^m)(x)$ et on obtient $x = (gf)^{m-n}(x) \in B$ ce qui est contradictoire. Un argument évident de cardinalité permet alors d'affirmer que $\text{codom}(\text{Ex}(f, g)) = C$. $\blacktriangleright$

Théorème 1.4 (Associativité de l'exécution)



$$\text{Ex}_n(\text{Ex}_n(f, g), h) = \text{Ex}_n(f, g + h) = \text{Ex}_n(\text{Ex}_n(f, h), g)$$

et ainsi

$$\text{Ex}(\text{Ex}(f, g), h) = \text{Ex}(f, g + h) = \text{Ex}(\text{Ex}(f, h), g)$$

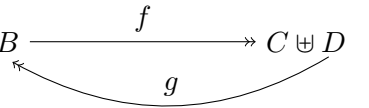
Preuve Soit $p \in \text{dom}(\text{Ex}_n(f, g + h))$, il existe un $m \leq n \in \mathbb{N}$ tel que

$$\begin{aligned}
 \text{Ex}_n(f, g + h)(p) &= f((g + h)f)^m(p) \\
 &= (f(gf)^{i_1})h \dots h(f(gf)^{i_k})(p) \text{ avec } i_1 + \dots + i_k + k - 1 = m \\
 &= (\text{Ex}_n(f, g)h\text{Ex}_n(f, g) \dots h\text{Ex}_n(f, g))(p) \\
 &= (\text{Ex}_n(f, g)(h\text{Ex}_n(f, g))^{k-1})(p) \\
 &= \text{Ex}_n(\text{Ex}_n(f, g), h)(p)
 \end{aligned}$$

En vertu de la commutativité de $+$ on obtient l'autre égalité. Ces égalités se transmettent naturellement à Ex . $\blacktriangleright$

Ce théorème a une signification très forte : c'est une version complètement localisée de la propriété de Church-Rosser. En effet, on verra dans la suite que des résultats de confluence forte peuvent être directement déduits de ce théorème.

La propriété suivante nous sera utile par la suite.

Propriété 1.5 Soient f, g et h des injections partielles telles que $A \uplus B \xrightarrow{f} C \uplus D$ 

et $\text{dom}(h) \cap \text{dom}(f) = \text{codom}(h) \cap \text{dom}(g) = \emptyset$.

On a $h + \text{Ex}(f, g) = \text{Ex}(f + h, g)$.

Preuve Cette propriété est directement issue du fait que $(f + h)g(f + h) = fgf$. $\blacktriangleright$

1.2.4 w -permutations et Ex -composition

Définition 1.6 On appelle w -permutation une permutation partielle involutive de domaine fini.

Une w -permutation est donc un produit de cycles disjoints de longueurs au plus 2.

Soit σ et τ des w -permutations disjointes et soit f une injection partielle avec $\text{dom}(f) \subseteq \text{dom}(\sigma)$ et $\text{codom}(f) \subseteq \text{dom}(\tau)$. On appelle Ex_0 -composition de σ et τ selon f la permutation partielle

$$\sigma \overset{f}{\rightsquigarrow}_0 \tau = \text{Ex}(\sigma + \tau, f + f^*)$$

La figure 1.3 donne une représentation de cette composition.

Propriété 1.7 $\sigma \overset{f}{\rightsquigarrow}_0 \tau$ est une w -permutation.

Preuve Soit x un élément de $\text{dom}(\sigma \overset{f}{\rightsquigarrow}_0 \tau)$, il existe un n tel que

$$(\sigma \overset{f}{\rightsquigarrow}_0 \tau)(x) = (\sigma + \tau)[(f + f^*)(\sigma + \tau)]^n(x)$$

Notons que $(\sigma + \tau)^* = \sigma + \tau$ et $(f + f^*)^* = f + f^*$, et ainsi on a $((\sigma + \tau)[(f + f^*)(\sigma + \tau)]^n)^* = [(\sigma + \tau)(f + f^*)]^n(\sigma + \tau) = (\sigma + \tau)[(f + f^*)(\sigma + \tau)]^n$. Donc $(\sigma \overset{f}{\rightsquigarrow}_0 \tau)^2(x) = x$. $\blacktriangleright$

On a naturellement $\text{dom}(\sigma \overset{f}{\rightsquigarrow}_0 \tau) = (\text{dom}(\sigma) - \text{dom}(f)) \uplus (\text{dom}(\tau) - \text{codom}(f))$. On peut se demander ce qui se passe pour les éléments de $(\text{dom}(\sigma) \cap \text{dom}(f)) \uplus (\text{dom}(\tau) \cap \text{codom}(f))$ par lesquels ne passe aucun calcul de **Ex**. Pour un tel élément dans $\text{dom}(\sigma) \cap \text{dom}(f)$, il existe nécessairement un n tel que $(f^* \tau f \sigma)^n(x) = x$. On obtient ainsi une sorte d'orbite $O_x = \{(f^* \tau f \sigma)^i(x), i \in \mathbb{N}\}$. A cette orbite, on peut associer une orbite duale pour τ par $O'_x = \{(f \sigma f^* \tau)((f \sigma)(x)), i \in \mathbb{N}\}$. On passe de O_x à O'_x bijectivement par application de $f \sigma$. On appelle *orbite double* l'ensemble formé par $O_x \cup O'_x$. On note $\mathbb{O}(\sigma \overset{f}{\rightsquigarrow}_0 \tau)$ l'ensemble de ces doubles orbites. Soit $R = \{\min_{x \in O \cup O'} x, (O, O') \in \mathbb{O}(\sigma \overset{f}{\rightsquigarrow}_0 \tau)\}$. On définit la **Ex**-composition, noté $\sigma \overset{f}{\rightsquigarrow} \tau$, de domaine $\text{dom}(\sigma \overset{f}{\rightsquigarrow}_0 \tau) \uplus R$ et telle que $\forall r \in R, (\sigma \overset{f}{\rightsquigarrow} \tau)(r) = r$.

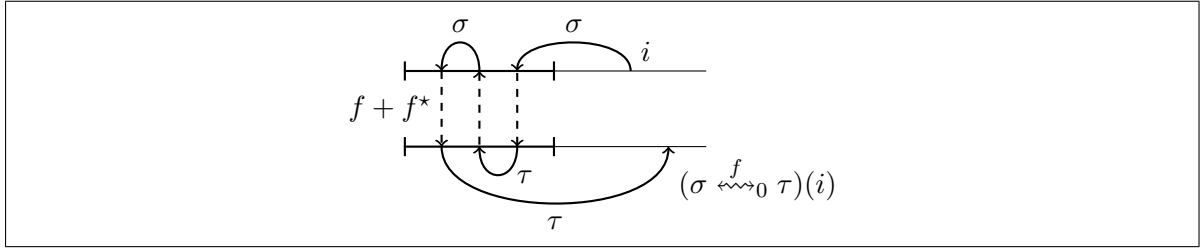
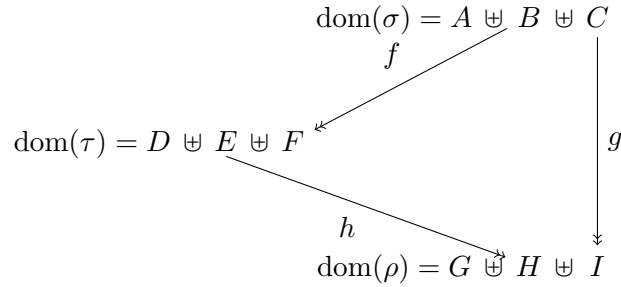


Figure 1.3: Représentation de la **Ex**₀-composition $\sigma \overset{f}{\rightsquigarrow}_0 \tau$

Nous donnons maintenant une propriété déduite du théorème 1.4, énonçant une sorte d'associativité pour la **Ex**-composition.

Propriété 1.8 Soient σ, τ, ρ des w -permutations disjointes deux à deux et



On a $\sigma \overset{f+g}{\rightsquigarrow} (\tau \overset{h}{\rightsquigarrow} \rho) = (\sigma \overset{f}{\rightsquigarrow} \tau) \overset{g+h}{\rightsquigarrow} \rho = (\sigma \overset{g}{\rightsquigarrow} \rho) \overset{f+h^*}{\rightsquigarrow} \tau$. Notamment lorsque $h = 0$ on a $\sigma \overset{f+g}{\rightsquigarrow} (\tau + \rho) = (\sigma \overset{f}{\rightsquigarrow} \tau) \overset{g}{\rightsquigarrow} \rho = (\sigma \overset{g}{\rightsquigarrow} \rho) \overset{f}{\rightsquigarrow} \tau$.

Preuve Il s'agit en fait de montrer d'une part la propriété pour la **Ex**₀-composition et d'autre part un résultat sur les orbites doubles.

On a $\sigma \overset{f+g}{\rightsquigarrow}_0 (\tau \overset{h}{\rightsquigarrow}_0 \rho) = \text{Ex}(\sigma + \text{Ex}(\tau + \rho, h + h^*), f + f^* + g + g^*) = \text{Ex}(\text{Ex}(\sigma + \tau + \rho, h + h^*), f + f^* + g + g^*)$ par la propriété 1.5 puis $= \text{Ex}(\sigma + \tau + \rho, f + f^* + g + g^* + h + h^*)$ par le théorème 1.4. On montre de la même manière que les autres compositions aboutissent à cette expression.

Au niveau des orbites doubles, posons

$$\mathbb{O}_3 = \mathbb{O}(\sigma \overset{f+g}{\rightsquigarrow}_0 (\tau \overset{h}{\rightsquigarrow}_0 \rho)) - \mathbb{O}(\tau \overset{h}{\rightsquigarrow}_0 \rho) - \mathbb{O}(\sigma \overset{f}{\rightsquigarrow}_0 \tau) - \mathbb{O}(\tau \overset{h}{\rightsquigarrow}_0 \rho)$$

Pour conclure il nous suffit de montrer que les orbites doubles de $\mathbb{O}_3$ ne dépendent pas de l'ordre dans lequel on effectue la composition. Or, une telle orbite est engendrée par un élément x tel que

$$\begin{aligned} x &= ((f^\star + g^\star)(\tau \xrightarrow{h}_0 \rho)(f + g)\sigma)^n(x) \\ &= \left(\prod_{i=1}^n (f^\star + g^\star)(\tau + \rho)((h + h^\star)(\tau + \rho))^{k_i}(f + g)\sigma\right)(x) \\ &= \left(\prod_{i=1}^n F_i\right)(x) \end{aligned}$$

où les F_i sont des quatre formes suivantes :

$$g^\star \rho h \tau (h^\star \rho h \tau)^{k_i} f \sigma, f^\star \tau (h^\star \rho h \tau)^{k_i} f \sigma, g^\star \rho (h \tau h^\star \rho)^{k_i} g \sigma, \text{ et } f^\star \tau h^\star \rho (h \tau h^\star \rho)^{k_i} g \sigma$$

En particulier lorsque $k_i = 0$ on peut avoir la forme $g^\star \rho g \sigma$ et $f^\star \tau f \sigma$. Cela suffit à pouvoir regrouper $\sigma \xrightarrow{g}_0 \rho$ et $\sigma \xrightarrow{f}_0 \tau$, et donc à retrouver les expressions de doubles orbites pour d'autres ordres de composition. ◀

Chapitre 2

Réseaux d'interaction explicites

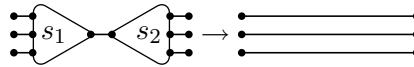
Une version courte de ce chapitre a fait l'objet de la publication [dF09].

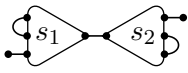
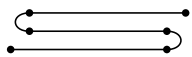
2.1 Motivations

Les réseaux d'interaction furent introduits par Lafont dans l'article [Laf90] comme un moyen d'extraire un modèle de calcul des réseaux de preuve de la logique linéaire multiplicative. Ils ont depuis été largement utilisés comme formalisme pour l'étude de l'implémentation de stratégies de réduction pour le λ -calcul, fournissant ainsi un cadre intuitif et visuel pour effectuer des substitutions explicites [Mac98, MP98, Lip03].

Même si la définition des réseaux d'interaction donnée au chapitre précédent suffit amplement pour pouvoir les utiliser, elle est toutefois trop limitée pour raisonner sur ceux-ci. Un des obstacles principaux provient du fait que leur nature est imprécise. Cette situation est similaire à celle des graphes : il est clair que nous ne pouvons les étudier en utilisant uniquement des dessins. De ce fait, on est amené à donner une définition précise d'un graphe, par exemple en tant que relation binaire ou en tant que (multi)ensemble d'arêtes.

Le principal obstacle à une telle définition réside dans la réduction. A titre d'exemple, considérons la règle



Peut-on l'appliquer sur le réseau d'interaction  ? Si nous sommes rigoureux, la partie gauche de la règle n'est pas exactement contenu dans le réseau dans la mesure où  n'est pas contenu dans . Peut-être devons nous considérer ce dernier fil comme composé de trois fils plus petits reposant sur deux ports temporaires : , et le réseau complet après réduction deviendrait . Mais à ce moment là, pour revenir à un vrai réseau d'interaction il faudrait concaténer tous ces fils et effacer les ports temporaires, ce qui donnerait le réseau . On fera référence ici à ce processus de concaténation des fils sous le terme de *fusion de ports*.

Il existe de nombreux travaux donnant des définitions des réseaux d'interaction et permettant une description rigoureuse de la réduction. Cependant, ils ont tous un point commun : ils traitent de manière implicite ou externe la fusion de ports. Dans l'article fondateur [Laf90], on trouve une définition des réseaux comme des ensembles de termes contenant au plus deux

occurrences de chaque variable, une variable représentant ainsi un fil. Cette présentation est raffiné et étendu dans l'article [FM99]. Dans ce cadre, c'est une relation d'équivalence sur les variables qui effectue la fusion de ports. Dans l'article [Pin00], une machine d'évaluation des réseaux est donnée, le calcul de la relation d'équivalence y alors présent sous forme d'une série d'étapes. Une approche rigoureuse partageant de nombreux outils avec la notre est donnée dans [Vau07], la fusion de ports y est faite à travers un algorithme externe de réécriture sur les ports.

Ainsi, on est amené à se poser la question suivante : pouvons-nous donner une définition des réseaux d'interaction dont on puisse extraire une description simple et rigoureuse de la réduction, y compris du processus de fusion de ports, tout en étant capable de retrouver les résultats standards, comme la confluence forte ? Le but de ce chapitre est de répondre à cette question.

Notre proposition repose sur l'observation suivante. Quand on branche la partie droite d'une règle dans un réseau, de nouveau fils sont définis par un processus d'aller-retour entre le réseau original et cette partie droite. Ce genre d'interaction sont prépondérantes dans la Géométrie de l'Interaction ou la *sémantique des jeux* [AJM94, HO00]. L'absence de typage dans les réseaux d'interaction fait de la GdI un possible cadre pour les exprimer. Pour cela on peut exprimer un réseau d'interaction par une permutation partielle et utiliser une composition basée sur la *formule d'exécution*. Une telle présentation des réseaux de preuves multiplicatifs a été effectuée par Girard dans l'article [Gir87]. Si l'on raisonne sur les actions dont on a fondamentalement besoin pour travailler avec des réseaux d'interaction, il est clair que l'on peut distinguer une *action des fils* consistant à aller d'un port à un autre le long d'un fil, et une *action des cellules* consistant à aller d'un port de cellule à un autre à l'intérieur d'une cellule. Ces deux actions amènent à décrire un réseau par un couple de permutations partielles. On peut alors se demander s'il est possible de combiner fidèlement celles-ci en une unique permutation, une solution à cette question est ce que l'on pourrait appeler une *Géométrie de l'Interaction*.

Le problème de la fusion de ports n'est pas inhérent aux réseaux d'interaction et on le retrouve dans d'autres contextes. La réécriture de diagrammes [Laf03] utilise une catégorie sous-jacente permettant mathématiquement le *redressement* des fils. Mais cela a un coût : la présentation repose sur un typage et est orientée verticalement, perdant ainsi la facilité de définition des réseaux d'interaction pour décrire des programmes. Un autre travail relié à ce problème est la présentation des réseaux multiplicatifs par Hughes dans [Hug05] où l'auteur présente les réseaux comme des fonctions munies d'une composition basée sur une construction catégorique associée aux catégories monoïdales tracées [JSV96] qui a été utilisée pour analyser la GdI [AJ92, HS06]. Une grande partie de notre cadre peut, en fait, être vu comme un cadre particulier de cette même construction générale. En effet, nous utiliserons les mêmes outils que dans ces sémantiques, cependant notre spécialisation à des injections partielles d'entiers permet de travailler sur la syntaxe et de rester dans un monde complètement non typé.

2.2 Statique

Fixons un ensemble dénombrable $\mathcal{S}$, dont les éléments sont appelés *symboles* et une fonction $\alpha : \mathcal{S} \rightarrow \mathbb{N}$, appelée *arité*. On définit les réseaux au-dessus de $\mathbb{N}$ et dans ce contexte un entier sera appelé un *port*.

Définition 2.1 *Un réseau d'interaction est un couple $R = (\sigma_w, \sigma_c)$ où*

- σ_w est une w -permutation. On note $P_l(R)$ ses points fixes et $P(R)$ les autres éléments de son domaine.
- σ_c est une permutation partielle de $P(R)$ à orbites pointées et étiquetées par $\mathcal{S}$ en sorte que $\forall o \in \text{Orbs}(\sigma_c), |o| = \alpha(l(o)) + 1$ où l est la fonction d'étiquetage.

Les éléments de $P_l(R)$ sont appelés *boucles* et les orbites restantes de σ_w , qui sont nécessairement d'ordre 2, sont appelées *fil*s. Le domaine de σ_w est appelé *support* du réseau. On note $P_c(R) = \text{dom}(\sigma_c)$, dont les éléments sont appelés *ports de cellules*, et $P_f(R) = P(R) - P_c(R)$, dont les éléments sont appelés *ports libres*.

Une orbite de σ_c est appelée une *cellule*. On note pal la fonction de pointage de σ_w . Soit c une cellule, $\text{pal}(c)$ est son *port principal* et pour tout $0 < i < |c|$ l'élément $(\sigma_c^i \circ \text{pal})(c)$ est son *ième port auxiliaire*.

Remarquons qu'une conséquence directe de cette définition est qu'un port est présent dans exactement un fil et au plus une cellule.

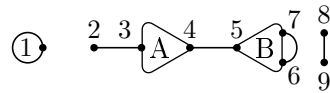
2.2.1 Représentation

Les réseaux admettent une représentation graphique très naturelle. On dessine une cellule de symbole A comme un triangle  où le port principal est le point au sommet et les ports auxiliaires sont alignés le long du côté opposé. On dessine les ports libres comme des points. Pour finir le dessin on ajoute des lignes entre deux ports reliés par un fil, et dessine des cercles pour les boucles.

A titre d'exemple, considérons le réseau $R = (\sigma_w, \sigma_c)$ où

$$\sigma_w = (1)(2\ 3)(4\ 5)(6\ 7)(8\ 9) \text{ et } \sigma_c = (\overset{\bullet}{4}\ 3)_A(\overset{\bullet}{5}\ 6\ 7)_B$$

les permutations étant données par leur décomposition en cycles et $(\overset{\bullet}{c_1}\ c_2\ \dots\ c_n)_S$ représente une cellule de point c_1 et de symbole S . Ce réseau aura la représentation suivante :



2.2.2 Morphismes de réseaux et renommage

Définition 2.2 Soient $R = (\sigma_w, \sigma_c)$ et $S = (\tau_w, \tau_c)$ des réseaux d'interaction. La fonction $f : \text{dom}(\sigma_w) \mapsto \text{dom}(\tau_w)$ est un morphisme de R dans S si et seulement si

$$f \circ \sigma_w = \tau_w \circ f, \quad f(P_c(R)) \subseteq P_c(R'),$$

$$\forall p \in P_c(R), (f \circ \sigma_c)(p) = (\tau_c \circ f)(p),$$

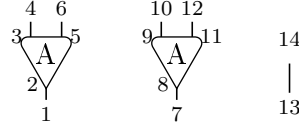
et $\forall o \in \text{Orbs}(\sigma_c)$ on a $(f \circ \text{pal})(o) = (\text{pal} \circ f)(o)$ et $l(o) = (l \circ f)(o)$.

Lorsque f est l'identité sur $P_f(R)$ on dit que le morphisme est interne.

A titre d'exemple considérons le réseau :

$$R = ((1\ 2)(3\ 4)(5\ 6)(7\ 8)(9\ 10)(11\ 12)(13\ 14), (\overset{\bullet}{2}\ 3\ 5)_A(\overset{\bullet}{8}\ 9\ 11)_A)$$

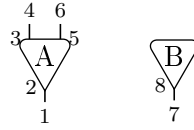
ayant la représentation :



et le réseau :

$$S = ((1\ 2)(3\ 4)(5\ 6)(7\ 8), (\overset{\bullet}{2}\ 3\ 5)_A(\overset{\bullet}{8})_B)$$

ayant la représentation :



Soit f l'application définie par $f(1) = f(7) = f(13) = 1$, $f(2) = f(8) = f(14) = 2$, $f(3) = f(9) = 3$, $f(5) = f(11) = 5$, $f(4) = f(10) = 4$ et $f(6) = f(12) = 6$. C'est un morphisme de R dans S .

Notons que l'égalité $f \circ \sigma_w = \tau_w \circ f$ est apparemment très forte, mais elle peut être déduite d'une simple inclusion, au sens des graphes d'applications, $f \circ \sigma_w \subseteq \tau_w \circ f$. En effet, soit (p, p') dans le graphe de $\tau_w \circ f$, on peut calculer $(f \circ \sigma_w)(p)$ qui par l'inclusion n'a pas d'autre choix que d'être égal à p' .

Détaillons cette définition. On rappelle que si σ et τ sont deux permutations et f une fonction du domaine de la première dans celui de la seconde, alors l'équation $f \circ \sigma = \tau \circ f$ a pour conséquence que toute orbite $o \in \text{Orbs}(\sigma)$ est envoyée sur une orbite $f(o) \in \text{Orbs}(\tau)$ telle que $|f(o)|$ soit un diviseur de $|o|$.

Dans le cas présent, une boucle est envoyée sur une boucle, un fil soit sur une boucle soit sur un fil, et une cellule sur une autre cellule. Les deux dernières équations de la définition assurent que le port principal d'une cellule est envoyé sur le port principal de la cellule image, et que le symbole est le même. La condition sur l'arité assure donc que chaque port est envoyé sur un port de même nature (principal, ou *i*ème port auxiliaire). De plus, seul un fil liant des ports libres peut être envoyé sur une boucle ou un fil reliant des ports quelconques. En effet, dès qu'un fil relie un port de cellule la troisième condition sur le morphisme force à l'envoyer sur un fil reliant également un port de cellule.

Avec ces remarques en tête, il est naturel d'appeler *renommage* (resp. *renommage interne*) un isomorphisme (resp. un isomorphisme interne). Une classe d'isomorphisme capture les réseaux d'interactions tels qu'ils seraient dessinés sur un papier magique où les ports peuvent bouger librement. Une classe d'isomorphisme interne correspond, quant à elle, aux réseaux ainsi dessinés pour lesquels on a de surcroît donné des noms distincts aux ports libres, justifiant ainsi le nom d'*interne*. Ceci est une notion importante car le dessin  est identique à

. Alors que le dessin, avec ports libres étiquetés $\begin{array}{cc} c & d \\ a & b \end{array}$ est différent de $\begin{array}{cc} c & d \\ a & b \end{array}$.

En fait, dès qu'on voudra considérer les réseaux comme représentant des sortes de termes, il sera nécessaire de ne les considérer qu'à isomorphisme interne près. Les ports libres représentent alors les variables libres alors que les ports de cellules correspondent aux variables liées. Par exemple, le λ -terme $\lambda x.(x)y$ est évidemment égal à $\lambda z.(z)y$ mais pas à $\lambda x.(x)z$.

Remarque 2.3 *Étant donné que les réseaux sont à support fini, nous pouvons toujours considérer que deux réseaux sont à supports disjoints modulo un renommage.*

Définition-Propriété 2.4 Soit $f : R \rightarrow S$ un morphisme de réseaux d'interaction. On peut définir le réseau image de R par f , noté $f(R)$, dont le support est inclus dans le support de S , et tel que la co-restriction de f à celui-ci soit un renommage.

Preuve Notons $R = (\sigma_w, \sigma_c)$, $S = (\tau_w, \tau_c)$ et $E = \text{codom}(f)$. Posons $\tau'_w = \tau_w \upharpoonright_E$, $\tau'_c = \tau_c \upharpoonright_E$, et $f(R) = (\tau'_w, \tau'_c)$. Le fait que cela soit un réseau est directement issu de la définition d'un morphisme. Le point essentiel est que cette définition entraîne que pour passer de τ_w (resp. τ_c) à τ'_w (resp. τ'_c) on n'a divisé aucune orbite. $\blacktriangleright$

2.3 Opérations sur les réseaux d'interaction

On donne ici les principaux outils et opérations qui vont être cruciaux dans notre définition de la réduction.

2.3.1 Recollement et découpage

Définition 2.5 Soient $R = (\sigma_w, \sigma_c)$ et $S = (\tau_w, \tau_c)$ deux réseaux à supports disjoints¹ et f une injection partielle de domaine inclus dans $P_f(R)$ et de co-domaine inclus dans $P_f(S)$. On appelle recollement de R et S le long de f le réseau $R \overset{f}{\rightsquigarrow} S = (\sigma_w \overset{f}{\rightsquigarrow} \tau_w, \sigma_c + \tau_c)$.

De la définition on déduit directement les faits suivants :

$$\begin{aligned} P(R \overset{f}{\rightsquigarrow} S) &= (P(R) - \text{dom}(f)) \uplus (P(S) - \text{codom}(f)) \\ P_c(R \overset{f}{\rightsquigarrow} S) &= P_c(R) \uplus P_c(S) \\ P_f(R \overset{f}{\rightsquigarrow} S) &= (P_f(R) - \text{dom}(f)) \uplus (P_f(S) - \text{codom}(f)) \\ R \overset{f}{\rightsquigarrow} S &= S \overset{f^*}{\rightsquigarrow} R \end{aligned}$$

Dans le cas d'un recollement le long de $f = 0$ on a $R \overset{0}{\rightsquigarrow} S = (\sigma_w + \tau_w, \sigma_c + \tau_c)$, on note ce recollement $R + S$, cela correspond à la notion usuelle de composition parallèle.

La figure 2.1 donne une représentation du recollement.

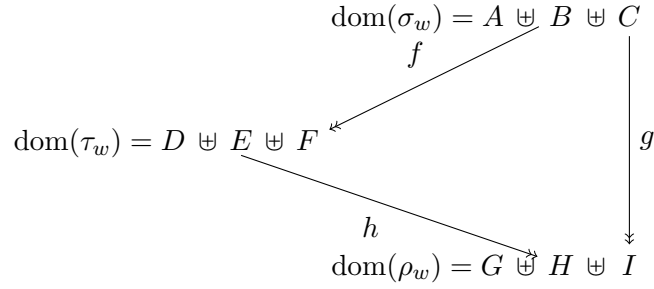
Propriété 2.6 Si $R = R \overset{f}{\rightsquigarrow} S$ alors $f = 0$ et $S = \mathbf{0} = (0, 0)$. Si $\mathbf{0} = R \overset{f}{\rightsquigarrow} S$ alors $f = 0$ et $R = S = \mathbf{0}$.

Preuve Les deux affirmations étant similaires, on ne va prouver que la première. Celle-ci est une conséquence directe des faits précédents : S devrait avoir ni cellules, ni ports libres et ni boucles. Le seul réseau ayant cette propriété est bien entendu le réseau vide $\mathbf{0}$. $\blacktriangleright$

On a une sorte d'associativité pour le recollement :

Propriété 2.7 Soient $R = (\sigma_w, \sigma_c)$, $S = (\tau_w, \tau_c)$ et $T = (\rho_w, \rho_c)$ des réseaux à supports disjoints deux à deux et

¹Ce qui n'est pas une perte de généralité en vertu de la remarque 2.3.



$$\text{On a } R \xrightarrow{f+g} (S \xrightarrow{h} T) = (R \xrightarrow{f} S) \xrightarrow{g+h} T = (R \xrightarrow{g} T) \xrightarrow{f+h} S.$$

Preuve La partie de cette égalité concernant les fils est une restriction du corollaire 1.8 et la partie concernant les cellules est issue de l'associativité de $+$. $\blacktriangleleft$

En pratique, c'est plus souvent le corollaire suivant que l'on utilisera :

Corollaire 2.8 *Supposons donnée une décomposition $R_0 = R \xrightarrow{f} (S \xrightarrow{g} T)$ alors il existe f_S et f_T telles que $R_0 = (R \xrightarrow{f_S} S) \xrightarrow{g+f_T} T$.*

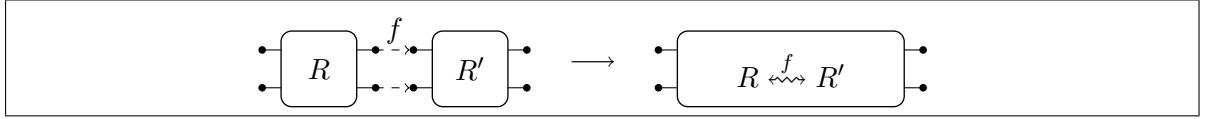


Figure 2.1: Représentation du recollement de deux réseaux d'interactions.

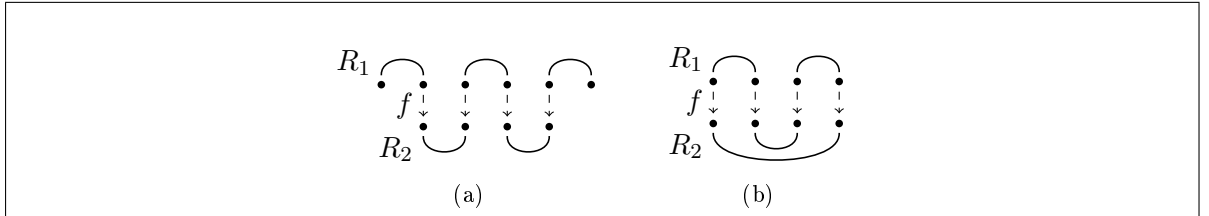


Figure 2.2: Représentation de deux découpages particuliers : (a) un découpage d'un simple fil et (b) un découpage d'une boucle

On peut définir une notion duale du recollement : le *découpage* consistant à extraire un sous-réseau d'un réseau d'interaction.

Définition 2.9 *Soit R un réseau, on appelle découpage de R un triplet (R_1, f, R_2) où $R = R_1 \xrightarrow{f} R_2$. Un réseau R' apparaissant dans un découpage de R est appelé un sous-réseau de R , ce que l'on note $R' \subseteq R$.*

La figure 2.2 présentent des exemples de découpage. Le fait que l'on puisse découper plusieurs fois un fil ou que l'on puisse diviser une boucle en un nombre quelconque de fils sont autant d'indices de la complexité inhérente à cette notion de sous-réseau

Propriété 2.10 *La relation $\subseteq$ est un ordre sur les réseaux.*

Preuve La relation $\subseteq$ est **réflexive** : $R = R \overset{0}{\rightsquigarrow} \mathbf{0}$ et ainsi, $R \subseteq R$.

Elle est **antisymétrique** : soient R_1 et R_2 des réseaux tels que $R_1 \subseteq R_2$ et $R_2 \subseteq R_1$. On a $R_1 = R_2 \overset{f}{\rightsquigarrow} R'_2$ et $R_2 = R_1 \overset{g}{\rightsquigarrow} R'_1$.

Ainsi $R_2 = (R_2 \overset{f}{\rightsquigarrow} R'_2) \overset{g}{\rightsquigarrow} R'_1$. En appliquant le corollaire 2.8 on obtient $R_2 = R_2 \overset{f_1}{\rightsquigarrow} (R'_2 \overset{g+f_2}{\rightsquigarrow} R'_1)$ et en appliquant la propriété 2.6, deux fois de suite, on obtient $R'_2 = R'_1 = \mathbf{0}$. On a donc $R_1 = R_2$.

Pour finir, la relation est **transitive** : soient $R \subseteq S \subseteq T$, alors on a $S = R \overset{f}{\rightsquigarrow} R'$ et $T = S \overset{g}{\rightsquigarrow} S'$, ainsi $T = (R \overset{f}{\rightsquigarrow} R') \overset{g}{\rightsquigarrow} S'$. En appliquant le corollaire 2.8 on obtient $T = R \overset{f_1}{\rightsquigarrow} (R' \overset{g+f_2}{\rightsquigarrow} S')$, ce qui implique que $R \subseteq T$. $\blacktriangleright$

2.3.2 Interfaces et contextes

Pour définir la réduction à l'aide de la relation de sous-réseau, il serait plus simple que l'on puisse référer implicitement à la fonction d'identification utilisée par le recollement. En guise d'intuition, considérons des contextes de termes ayant un nombre quelconque de trous, pour substituer complètement un tel contexte on peut se donner une fonction des trous dans les termes et les remplir ainsi. Cependant, une définition plus naturelle serait d'attribuer un numéro à chaque trou, par exemple son ordre d'apparence de gauche à droite, et puis de se donner une liste de termes. La substitution reviendrait alors à remplir le i ème trou par le i ème terme de la liste. La définition qui va suivre est une transposition de ces intuitions dans le cadre des réseaux d'interaction.

Définition 2.11 On appelle interface d'un réseau R un sous-ensemble de $P_f(R)$ ordonné linéairement, la longueur de la chaîne étant appelée taille. On dit, alors, que R contient l'interface I , noté $I \subset R$. Une interface est canonique si elle contient tous les ports libres d'un réseau.

Soit I et I' des interfaces disjointes d'un même réseau, on note II' l'union de ces sous-ensembles ordonnée par la concaténation des deux chaînes d'ordre. Plus précisément : $x \leq_{II'} y \iff x \leq_I y$ ou $x \leq_{I'} y$ ou $x \in I \wedge y \in I'$.

Soient I et I' deux interfaces de même taille, il existe une et une seule bijection de I dans I' préservant les ordres associés, on la note $\rho(I, I')$ et on l'appelle le raccord entre I et I' .

On appelle contexte un couple (R, I) où I est une interface contenue dans le réseau R , on le note R^I .

Soient R^I et $R'^{I'}$ deux contextes aux interfaces de même taille, on note

$$R^I \rightsquigarrow R'^{I'} = R \overset{\rho(I, I')}{\rightsquigarrow} R'$$

Dans la suite, lorsque l'on écrira $R^I \rightsquigarrow R'^{I'}$ on supposera implicitement que I et I' sont de même taille.

On peut maintenant énoncer littéralement la commutativité du recollement (la preuve étant triviale) :

Propriété 2.12 $R^I \rightsquigarrow R'^{I'} = R'^{I'} \rightsquigarrow R^I$

Le fait trivial suivant assure que tout recollement peut être vu comme un recollement de contextes.

Fait 2.13 Soit $R \xleftrightarrow{f} R'$ un recollement, il existe des interfaces $I \subset R$ et $I' \subset R'$ telles que $R \xleftrightarrow{f} R' = R^I \xleftrightarrow{\quad} R^{I'}$.

Corollaire 2.14 $R_1 \subseteq R \iff \exists I_1, R_2, I_2$ tels que $R = R_1^{I_1} \xleftrightarrow{\quad} R_2^{I_2}$

On peut énoncer à nouveau le corollaire 2.8 avec des interfaces :

Corollaire 2.15 Pour tous réseaux R, S, T et interfaces I, J, K, L , il existe des interfaces I', J', K', L' telles que

$$R^I \xleftrightarrow{\quad} (S^J \xleftrightarrow{\quad} T^K)^L = (R^{I'} \xleftrightarrow{\quad} S^{J'})^{L'} \xleftrightarrow{\quad} T^{K'}$$

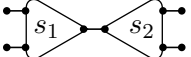
2.4 Dynamique

Définition 2.16 Soient s_1 et s_2 des symboles. On appelle règle d'interaction pour (s_1, s_2) un couple $(R_r^{I_r}, R_p^{I_p})$ où

$$R_r = \left(\begin{array}{c} (b \ c)(a_1 \ b_1) \dots (a_n \ b_n)(c_1 \ d_1) \dots (c_m \ d_m), \\ \bullet \\ (\dot{b} \ b_1 \dots b_n)_{s_1} (\dot{c} \ c_1 \dots c_m)_{s_2} \end{array} \right)$$

et I_r, I_p sont canoniques et de même taille.

Soit $\mathcal{R} = (R_r^{I_r}, R_p^{I_p})$ une règle, on appelle réduction par $\mathcal{R}$ la relation binaire $\xrightarrow{\mathcal{R}}$ sur les réseaux telle que pour tous renommages α et β , ainsi que pour tout réseau S tel que $S = R^I \xleftrightarrow{\quad} \alpha(R_r)^{\alpha(I_r)}$ on aie $S \xrightarrow{\mathcal{R}} S'$ où $S' = R^I \xleftrightarrow{\quad} \beta(R_p)^{\beta(I_p)}$.

Le réseau R_r a la représentation  . Remarquons que la réduction est définie dès que le réseau contient un renommé du redex R_r . Cette réduction donne l'illusion d'être non-déterministe mais il ne s'agit que de l'expansion d'une réduction déterministe pour tenir compte de tous les renommages possibles.

Propriété 2.17 Soient R un réseau et $\mathcal{R}_1, \mathcal{R}_2$ deux règles d'interaction applicable sur R sur des redex distincts tels que $R_1 \xleftarrow{\mathcal{R}_1} R \xrightarrow{\mathcal{R}_2} R_2$ et que l'intersection du support de R_1 et du support de R_2 soit contenu dans le support de R .

Il existe un réseau R' tel que $R_1 \xrightarrow{\mathcal{R}_2} R' \xleftarrow{\mathcal{R}_1} R_2$.

Preuve Pour $i = 1, 2$, on pose $\mathcal{R}_i = (R_{r,i}^{I_{r,i}}, R_{p,i}^{I_{p,i}})$. La forme des redex permet d'affirmer que s'ils sont distincts alors ils sont disjoints. Comme R contient à la fois le redex $\alpha_1(R_{r,1})$ et le redex $\alpha_2(R_{r,2})$, on peut en déduire que $\alpha_1(R_{r,1}) + \alpha_2(R_{r,2}) \subseteq R$. Plus précisément, on a

$$R = (\alpha_1(R_{r,1}) + \alpha_2(R_{r,2}))^{\alpha_1(I_{r,1})\alpha_2(I_{r,2})} \xleftrightarrow{\quad} R_0^I$$

On obtient donc

$$R_1 = (\beta_1(R_{p,1}) + \alpha_2(R_{r,2}))^{\beta_1(I_{p,1})\alpha_2(I_{r,2})} \xleftrightarrow{\quad} R_0^I$$

pour un renommage β_1 , et on a une expression similaire pour R_2 . Il est direct de vérifier que le réseau

$$R' = (\beta_1(R_{p,1}) + \beta_2(R_{p,2}))^{\beta_1(I_{p,1})\beta_2(I_{p,2})} \xleftrightarrow{\quad} R_0^I$$

satisfait la conclusion à l'aide la propriété 2.7. L'existence de ce réseau repose sur la disjonction des $\beta_i(R_{p,i})$ qui est une conséquence sur l'hypothèse à propos des ports contenus à la fois dans R_1 et R_2 . ◀

Corollaire 2.18 *Soit $\mathcal{B}$ un ensemble de règles tel que pour toute paire de symboles il existe au plus une règle sur ceux-ci dans $\mathcal{B}$. Un tel ensemble sera appelé une bibliothèque. La réduction $\xrightarrow{\mathcal{L}} = \bigcup_{R \in \mathcal{L}} \xrightarrow{R}$ est fortement confluente à un renommage près.*

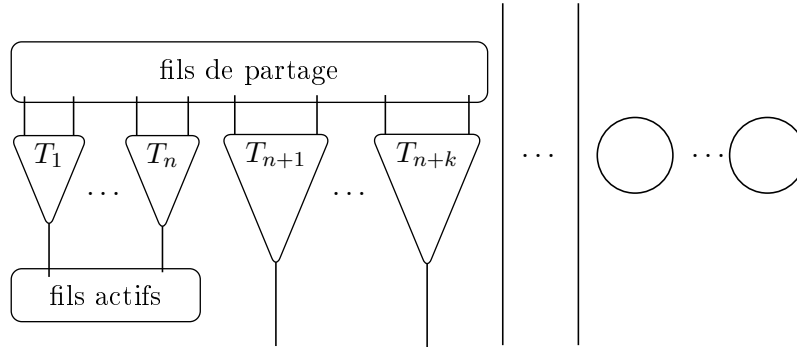
L'expression à un renommage près signifie que l'on peut avoir à effectuer un renommage d'un des réseaux d'une paire critique avant de pouvoir les rejoindre. Ceci est dû à la condition de disjonction pour appliquer la propriété 2.17.

2.5 Morphologie générale des réseaux

Soit R un réseau d'interaction. On peut classer les fils de R en quatre catégories :

1. les fils *actifs* reliant deux ports principaux : ces fils représentent les redex et la mise en communication immédiate de deux ressources, dont l'interaction est la réduction.
2. les fils *d'agrégation* reliant un port principal à un port auxiliaire : la cellule ainsi liée par son port principal est *inhibée*, elle ne peut pas réagir avant que l'autre cellule ne se réduise. Cela correspond bien à la notion de possession ou d'agrégation.
3. les fils *de partage* reliant deux ports auxiliaires.
4. les fils *libres* reliant au moins un port libre.

La forme générale d'un réseau d'interaction est la suivante :



Les T_i représentent des arbres de cellules reliées par des fils d'agrégation.

Notons que quelque soit la bibliothèque considérée, un réseau sans fil actif est en forme normale.

Définition 2.19 *Soit R un réseau d'interaction, si R n'a pas de fils actifs on dit que R est en forme normale absolue.*

2.6 Les réseaux d'interaction sont le Ex-effondrement des réseaux Axiome/Coupure

On introduit maintenant une notion de réseaux se situant entre les réseaux d'interaction et les réseaux de preuve de logique linéaire multiplicative.

Lorsque l'on connecte deux réseaux d'interaction, par identification de ports libres, il se produit un processus complexe de simplification des fils. C'est ce que nous avons pu voir avec l'utilisation de la Ex-composition. En revanche, pour connecter deux réseaux de preuve, il n'est

plus possible de directement opérer cette identification, ici entre conclusions, et l'on est obligé de rajouter un lieu entre les deux conclusions : *une coupure*. Cette composition est beaucoup plus simple, puisqu'elle ne fait que rajouter du contenu. Une étape supplémentaire est par contre nécessaire pour *éliminer ces coupures*.

La notion de réseaux que l'on introduit ici va donc faire intervenir des fils de deux sortes : les axiomes et les coupures. Le but final étant d'expliciter le slogan

Les réseaux d'interaction sont un quotient des réseaux de preuve multiplicatif

faisant partie du folklore.

2.6.1 Définition et juxtaposition

Définition 2.20 Un réseaux Axiome/Coupure, réseau AC en abrégé, est un triplet $R = (\sigma_A, \sigma_C, \sigma_c)$ où :

- σ_A et σ_C sont des w -permutations satisfaisant $\text{dom}(\sigma_C) \subseteq \text{dom}(\sigma_A)$, σ_C sans point fixe et :

bord d'une coupure si $(a\ b)$ est une orbite de σ_C alors il existe $c \neq a$ et $d \neq b$ tels que $(c\ a)$ et $(b\ d)$ soient des orbites de σ_A

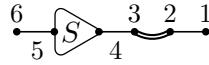
On note $P_l(R)$ les points fixes de σ_A et $P(R) = \text{dom}(\sigma_A) - \text{dom}(\sigma_C) - P_l(R)$.

- σ_c est une permutation partielle de $P(R)$ à orbites pointées et étiquetées par $\mathcal{S}$ en sorte que $\forall o \in \text{Orbs}(\sigma_c), |o| = \alpha(l(o))$ où l est la fonction d'étiquetage.

On adapte directement la représentation des réseaux d'interaction aux réseaux AC en affichant orbites de σ_C par des traits doubles. Par exemple le réseau AC $R = (\sigma_A, \sigma_C, \sigma_c)$ où

$$\sigma_A = (1\ 2)(3\ 4)(5\ 6), \sigma_C = (2\ 3), \sigma_c = (\overset{\bullet}{4}\ 5)_S$$

sera représenté par



On peut adapter la plupart des définitions précédentes à ces réseaux, plus particulièrement les définitions de ports libres, interfaces et contextes. L'avantage des réseaux AC est qu'ils admettent une composition extrêmement simple :

Définition 2.21 Soient $R^I = (\sigma_A, \sigma_C, \sigma_c)$ et $R^{I'} = (\tau_A, \tau_C, \tau_c)$ deux contextes sur des réseaux AC à supports disjoints, où $I = i_1 > \dots > i_n$ et $I' = i'_1 > \dots > i'_n$.

On appelle juxtaposition de R^I et $R^{I'}$ le réseau AC

$$R^I \leftrightarrow R^{I'} = (\sigma_A + \tau_A, \sigma_C + \tau_C + \rho(I, I') + \rho(I, I')^*, \sigma_c + \tau_c)$$

Du point de vue logique la juxtaposition est une coupure généralisée, en effet, son interprétation en termes de permutations correspond exactement à la définition de la coupure que l'on peut trouver dans [Gir87].

2.6.2 Ex-effondrement

On donne ici une variante restreinte de la procédure d'élimination des coupures de la logique.

Définition-Propriété 2.22 Soit $R = (\sigma_A, \sigma_C, \sigma_c)$ un réseau AC et f une injection partielle telle que $\text{dom}(\sigma_C) = \text{dom}(f)$ et $\text{codom}(f) \cap \text{dom}(\sigma_A) = \emptyset$.

Le couple $(\sigma_A \xrightarrow{f} f \circ \sigma_C \circ f^*, \sigma_c)$, est un réseau d'interaction.

Le résultat ne dépend pas de f et on l'appelle le Ex-effondrement de R , noté $\text{Ex}(R)$.

Pour que cette définition soit correcte, on a dû délocaliser σ_C vers un domaine disjoint de $\text{dom}(\sigma_A)$. Le Ex-effondrement procède donc ainsi : on remplace chaque chaîne maximale $a_1 \xrightarrow{\sigma_A} b_1 \xrightarrow{\sigma_C} a_2 \dots b_{n-1} \xrightarrow{\sigma_A} a_n$ par une chaîne

$$a_1 \xrightarrow{\sigma_A} b_1 \xrightarrow{f} f(b_1) \xrightarrow{f \circ \sigma_C \circ f^*} f(a_2) \xrightarrow{f^*} \dots b_{n-1} \xrightarrow{\sigma_A} a_n$$

et on calcule ensuite la Ex-composition pour obtenir $a_1 \xrightarrow{\sigma_A \xrightarrow{f} \sigma_C} a_n$.

Preuve Remarquons que pour que cela soit un réseau d'interaction, la seule propriété qui n'est pas issue de la définition des réseaux AC est le fait que $\sigma_A \xrightarrow{f} f \circ \sigma_C \circ f^*$ est une w-permutation, mais on le déduit de la propriété 1.7.

La remarque précédente assure que f n'a qu'un rôle artificiel dans la définition. En effet, à chaque fois que f est appliquée lors de la Ex-composition elle est suivie d'une application de son inverse. Plus généralement, pour des injections partielles $\sigma_1, \tau_1, \dots, \sigma_n, \tau_n$, on a $\sigma_n \circ \tau_n \circ \dots \circ \sigma_1 \circ \tau_1 = \sigma_n \circ f_n^* \circ f_n \circ \tau_n \circ g_n^* \circ g_n \circ \dots \circ g_1^* \circ g_1 \circ \sigma_1 \circ f_1^* \circ f_1 \circ \tau_1$ pour toutes injections partielles $f_1, g_1, \dots, f_n, g_n$ dès que $\text{dom}(f_i) \subseteq \text{codom}(\tau_i) \cap \text{dom}(\sigma_i)$ et $\text{dom}(g_i) \subseteq \text{codom}(\sigma_i) \cap \text{dom}(\tau_i)$ ².

✎

Propriété 2.23 Pour chaque réseau d'interaction R , il existe un unique réseau AC R' de la forme $(\sigma_A, 0, \sigma_c)$ tel que $\text{Ex}(R') = R$.

R' est dit sans coupures.

Preuve Si $R = (\tau_w, \tau_c)$ il suffit de poser $R' = (\tau_w, 0, \tau_c)$. L'unicité provient du fait que $\sigma \xrightarrow{0} 0 = \sigma$.

✎

Définition 2.24 Soient R et R' deux réseaux AC, on dit qu'ils sont Ex-équivalent, noté $R \rightsquigarrow R'$, lorsque $\text{Ex}(R) = \text{Ex}(R')$.

On a une correspondance directe entre la juxtaposition et le recollement :

Propriété 2.25 $\text{Ex}(R^I \leftrightarrow R'^{I'}) = \text{Ex}(R)^I \rightsquigarrow \text{Ex}(R')^{I'}$

Preuve Notons $R = (\sigma_A, \sigma_C, \sigma_c)$, $R' = (\tau_A, \tau_C, \tau_c)$.

Si on note f (resp. g) l'injection partielle utilisée dans le calcul de $\text{Ex}(R)$ (resp. $\text{Ex}(R')$), alors il est possible de trouver une injection partielle h telle que l'injection partielle utilisée pour effectuer le calcul de $\text{Ex}(R^I \leftrightarrow R'^{I'})$ soit $f + g + h$. De plus, on peut décomposer $h = i + i'$ afin que $h(\rho(I, I') + \rho(I, I')^*)h^* = i\rho(I, I')i^* + i'\rho(I, I')^*i'^*$.

² Remarquons qu'il s'agit là de la même technique que celle permettant d'exprimer un amalgame de groupes $*_A G_i$ comme quotient d'une action de A^n (cf. [Ser77]).

L'essentiel de la propriété revient donc à prouver que

$$(\sigma_A + \tau_A) \xrightarrow{f+g+i+i'} (f\sigma_C f^* + g\tau_C g^* + i\rho(I, I')i^* + i'\rho(I, I')^*i'^*) = \\ (\sigma_A \xrightarrow{f} f\sigma_C f^*) \xrightarrow{\rho(I, I')} (\tau_A \xrightarrow{g} g\tau_C g^*)$$

Cette égalité se déduit de manière similaire à la propriété 1.8, la présence des injections partielles f, g, i et i' ne rajoutant qu'une difficulté artificielle. $\blacktriangleleft$

On peut donc affirmer un slogan précis :

Les réseaux d'interaction sont le quotient des réseaux AC par $\approx$

2.7 La catégorie IN et l'approche DPO

2.7.1 Motivations pour l'approche DPO

L'approche DPO, pour *double push-out*, est une des approches algébriques de la réécriture de graphes. La réécriture de graphes étant très délicate à manipuler rigoureusement, il est essentiel d'avoir de bons outils algébriques pour la décrire.

Nous décrivons brièvement et grossièrement cette approche. On considère comme règle un diagramme $R \leftarrow I \rightarrow S$ dans **Graph**, la catégorie des graphes. Le graphe I correspond à une sorte d'interface commune entre R et S . Dès qu'on a un morphisme $R \rightarrow G$ ayant de bonnes propriétés, on dit alors que G contient le motif de la règle, on peut construire dans **Graph**

un graphe G' tel que l'on ait une somme amalgamée (en anglais. pushout)

$$\begin{array}{ccc} R & \xleftarrow{\quad} & I \\ \downarrow & \text{s.a.} & \downarrow \\ G & \xleftarrow{\quad} & G' \end{array} .$$

Nous verrons la définition précise des sommes amalgamées dans la suite, informellement le graphe G' correspond alors au complément de R dans G vis-à-vis de l'interface I . On peut encore construire un pushout dans l'autre sens, de manière à obtenir le double carré :

$$\begin{array}{ccccc} R & \xleftarrow{\quad} & I & \xrightarrow{\quad} & S \\ \downarrow & & \downarrow & & \downarrow \\ & \text{s.a.} & & \text{s.a.} & \\ G & \xleftarrow{\quad} & G' & \xrightarrow{\quad} & G_r \end{array}$$

Le graphe G_r est alors appelé réduit de G par la règle. Il a été construit en prenant le complément G' et en remplaçant ce qui était l'emplacement de R par le graphe S , puis en appliquant une opération de collage le long de l'interface I .

L'intérêt de cette approche est que, bien que la définition précise soit très catégorique, elle est néanmoins très proche de notre intuition. Cette approche a été initiée dans l'article [EPS73]. Elle a depuis été très étudiée, et étendue à de nombreux autres formalismes.

L'intuition derrière la réécriture de graphe est directement transposable en termes de réseaux d'interaction. De plus, ces notions informelles de collage et de découpage sont dans notre formalisme des notions algébriques. Il est alors assez naturel d'essayer de présenter la réécriture de réseaux par une approche DPO, c'est le but de cette partie.

Notons qu'il existe une telle approche dans la littérature [Ban95], elle repose alors sur un plongement des réseaux dans les hypergraphes, suivi d'un plongement de ceci dans les graphes bipartites. Les notions y sont ainsi moins directement accessibles que dans notre présentation. Cependant, dans la mesure où l'on peut dans les deux cas récupérer la réduction originelle des réseaux, il est possible que les deux approches soient les mêmes.

2.7.2 Définition

Soit IN la catégorie dont les objets sont les réseaux d'interaction et dont les morphismes sont les morphismes de réseaux.

Dans ce paragraphe on note $R \xrightarrow{f} S$ pour dire que f est un morphisme de R dans S , injectif en tant que fonction sur les ports. De même on note $R \xrightarrow{\sim} S$ lorsque f est une bijection.

Pour tout morphisme $R \xrightarrow{f} S$ on note $\tilde{f}$ la co-restriction de f à $f(R)$, on a donc $R \xrightarrow{\tilde{f}} f(R)$. $\tilde{f}$ est alors un renommage.

Lemme 2.26 *Si $R \xrightarrow{f} S$, alors $\tilde{f}(R) \subseteq S$.*

Réciproquement, si $\alpha(R) \subseteq S$ pour un renommage α , alors $R \xrightarrow{\hat{\alpha}} S$ où $\hat{\alpha}$ est la co-extension de α au support de S .

En terme de diagrammes, cela revient à dire que lorsqu'on a un morphisme $R \xrightarrow{f} S$ alors on peut trouver un iso f' et un réseau $S' \subseteq S$ tel que le diagramme

$$\begin{array}{ccc} & f' & S' \\ & \nearrow & \searrow \\ R & \xrightarrow{f} & S \end{array} \quad \text{commute.}$$

Réciproquement, dès que l'on a le haut du diagramme, alors on peut trouver une injection f tel que le triangle soit commutatif.

Preuve La première implication est issue de la définition du réseau image $f(R)$. La seconde implication est triviale. $\blacktriangleright$

2.7.3 Sommes amalgamées dans Set

On rappelle la définition des sommes amalgamées.

Définition 2.27 *Soit $\mathcal{C}$ une catégorie. Un carré commutatif*

$$\begin{array}{ccc} & R & \\ f \swarrow & & \searrow f' \\ S & & S' \\ g \searrow & T & \swarrow g' \end{array}$$

est appelé somme

amalgamée (en anglais push-out) si pour tout autre carré

$$\begin{array}{ccc} & R & \\ f \swarrow & & \searrow f' \\ S & & S' \\ h \searrow & T' & \swarrow h' \end{array}$$

il existe un unique

$T \xrightarrow{u} T'$ *tel que $ug = h$ et $ug' = h'$.*

On note qu'un carré est une somme amalgamée en marquant s.a. en son centre.

Si on peut compléter tout diagramme
$$\begin{array}{ccc} & f & R \\ & \swarrow & \searrow \\ S & & S' \end{array} \begin{array}{c} f' \\ \\ \end{array}$$
 en une somme amalgamée, on dit

que $\mathbf{C}$ a toutes les sommes amalgamées.

Rappelons tout d'abord un résultat classique.

Fait 2.28 *Set a toutes les sommes amalgamées.*

Preuve Soit $F \xleftarrow{f} E \xrightarrow{f'} F'$. Posons $G = F \uplus F'$, π_1 et π_2 les projections de F et F' dans G , et soit $\sim$ la plus petite relation d'équivalence définie sur G et telle que

$$\forall z \in E, \pi_1 f(z) \sim \pi_2 f'(z)$$

On note g la surjection canonique associée à $\sim$.

Le carré
$$\begin{array}{ccc} & f & E \\ & \swarrow & \searrow \\ F & & F' \\ & \searrow & \swarrow \\ & g\pi_1 & G/\sim \\ & & g\pi_2 \end{array}$$
 est, par définition de $\sim$, commutatif. Montrons qu'il s'agit

d'une somme amalgamée.

Supposons donné un autre carré commutatif
$$\begin{array}{ccc} & f & E \\ & \swarrow & \searrow \\ F & & F' \\ & \searrow & \swarrow \\ & h & H \\ & & h' \end{array}$$
. La somme disjointe étant

un coproduit dans **Set**, on a un unique morphisme $G \xrightarrow{u} H$ tel que $u\pi_1 = h$ et $u\pi_2 = h'$. Or, soient $x, y \in G$ tel que $x \sim y$. Par définition on a $x = \pi_1 f(z)$ et $y = \pi_2 f'(z)$ pour un $z \in E$. Donc $u(x) = u\pi_1 f(z) = hf(z)$ et $u(y) = u\pi_2 f'(z) = h'f'(z)$, et par commutativité du carré on a $u(x) = u(y)$. Donc, u se factorise de manière unique en $G \xrightarrow{g} G/\sim \xrightarrow{\bar{u}} H$. $\blacktriangleright$

Corollaire 2.29 *Dans Set lorsqu'on complète un diagramme $F \xleftarrow{f} E \xrightarrow{f'} F'$ en une somme amalgamée, alors les morphismes fermant le carré sont également des injections.*

Preuve Dans la preuve précédente, si f et f' sont des injections, alors c'est également le cas de $g\pi_1$ et $g\pi_2$. En effet, si par exemple $g\pi_1(p) = g\pi_1(p')$ cela signifie que $\pi_1(p) \sim \pi_1(p')$. Pour avoir cette équivalence, soit on a utilisé la réflexivité ce qui entraîne $p = p'$ car π_1 est injectif, soit on a utilisé la transitivité et il existe $z, z' \in E$ tel que $f'(z) = f'(z')$, $f(z) = p$ et $f(z') = p'$. Par injectivité de f' on a donc $z = z'$ qui entraîne $p = p'$. $\blacktriangleright$

2.7.4 Sommes amalgamées dans IN

2.7.4.a Concrétisation de IN

Soit $\mathbf{IN} \xrightarrow{\mathbf{U}} \mathbf{Set}$ le foncteur d'oubli qui a un réseau associe son support. C'est naturellement un foncteur fidèle car la donnée d'un morphisme f de réseaux est exactement la donnée de $\mathbf{U}(f)$ ³. On va donc pouvoir se servir des sommes amalgamées dans **Set** pour les construire

³Par contre, il n'est pas plein car ce ne sont que certaines applications entre supports qui sont des morphismes.

dans $\mathbf{IN}$.

2.7.4.b Obstacles

Observons tout d'abord deux obstacles à la constructions de sommes amalgamées, tous deux liés à la présence des cellules. Il est raisonnable qu'il nous faille utiliser des cellules pour trouver des problèmes, car la catégorie des graphes a toutes les sommes amalgamées et les réseaux sans cellules sont des graphes particuliers.

Soient $R = ((1\ 2)(3\ 4), 0)$, $S = ((5\ 6), 0)$ et $S' = ((7\ 8), 0)$, $f : R \rightarrow S$ définie par $f(1) = f(3) = 5, f(2) = f(4) = 6$ et $f' : R \rightarrow S'$ définie par $f(1) = f(4) = 7, f(2) = f(3) = 9$. Si on construit la somme amalgamée dans $\mathbf{Set}$ on obtient une unique classe d'équivalence $5 \sim 6 \sim 7 \sim 8$. L'unique candidat pour définir l'image de cet élément le long d'un fil est donc lui même : on a construit une boucle. Cela pose un problème dès qu'un des ports de R est connecté à une cellule : il n'y a aucun moyen de définir l'image de cette cellule dans l'amalgame.

On va pouvoir éviter ce problème en imposant l'injectivité. Cependant, même avec cela on peut quand même avoir de mauvais cas de figure.

Soit $R = ((1\ 2), 0)$, $S = ((3\ 4), (\overset{\bullet}{3})_A)$, $S' = ((5\ 6)(7\ 8), (\overset{\bullet}{9}\ 7)_B)$, $f : R \rightarrow S$ définie par $f(1) = 3, f(2) = 4$ et $f' : R \rightarrow S'$ définie par $f(1) = 5, f(2) = 6$. Dans le quotient des supports on a quatre classes d'équivalence : $3 \sim 5, 4 \sim 6, 7$ et 8 . Le seul candidat pour la w-permutation des fils est alors $(3 \sim 5\ 4 \sim 6)(7\ 8)$ comme on le verra plus tard. Il n'est cependant pas possible de définir des morphismes vers l'amalgame, car $3 \sim 5$ devrait à la fois être un port principal d'une cellule de symbole A et un port auxiliaire d'une cellule de symbole B .

On en déduit la condition nécessaire, et comme on le verra suffisante, suivante :

Définition 2.30 Soit $S \xleftarrow{f} R \xrightarrow{f'} S'$ un diagramme dans $\mathbf{IN}$. On dit que f et f' sont conciliables si et seulement si pour tout $p \in P(R) - P_c(R)$ tel que $f(p) \in P_c(S)$ et $f'(p) \in P_c(S')$, alors $f(p)$ et $f'(p)$ appartiennent à des cellules ayant même symbole et ils ont même nature, i.e. ce sont tous les deux des ports principaux ou des ième ports auxiliaires.

Dans les applications la condition suivante, plus forte, sera en général assurée :

Définition 2.31 Soit $S \xleftarrow{f} R \xrightarrow{f'} S'$ un diagramme dans $\mathbf{IN}$.

On dit que f et f' sont superposables si et seulement si pour tout $p \in P(R) - P_c(R)$ on a

$$f(p) \in P_c(S) \iff f'(p) \notin P_c(S')$$

On dit que f et f' sont fortement non superposables si et seulement si pour tout $p \in P(R) - P_c(R)$ on a

$$f(p) \in P_c(S) \iff f'(p) \in P_c(S')$$

2.7.4.c Sommes amalgamées de morphismes conciliables

Théorème 2.32 Soit $S \xleftarrow{f} R \xrightarrow{f'} S'$ un diagramme dans $\mathbf{IN}$ où f et f' sont conciliables, alors

on peut le compléter en

$$\begin{array}{ccccc} & & R & & \\ f \swarrow & & & \searrow f' & \\ S & & & & S' \\ & \searrow s.a. & & \swarrow & \\ & T & & & \end{array}$$

Preuve Soit $S \xleftarrow{f} R \xrightarrow{f'} S'$ avec $R = (\sigma_w, \sigma_c)$, $S = (\tau_w, \tau_c)$ et $S' = (\tau'_w, \tau'_c)$. Dans toute cette preuve, on identifie un morphisme de réseau f à $\mathbf{U}(f)$. Cela est justifié par le fait que $\mathbf{U}$ soit fidèle.

α) En reprenant la preuve du fait 2.28 et le corollaire 2.29, on a une somme amalgamée dans **Set** :

$$\begin{array}{ccc} & \mathbf{U}(R) & \\ f \swarrow & & \searrow f' \\ \mathbf{U}(S) & \text{s.a.} & \mathbf{U}(S') \\ g\pi_1 \searrow & & \swarrow g\pi_2 \\ & G/\sim & \end{array}$$

Tout d'abord, $G/\sim$ étant fini, nous pouvons l'identifier à une partie de $\mathbb{N}$. C'est donc un support valide pour un réseau d'interaction.

β) Soit $p \in G/\sim$, supposons qu'il soit de la forme $p = g\pi_1(p_1)$ on pose $\rho_w(p) = g\pi_1\tau_w(p_1)$. On fait de même vis-à-vis de S' . Si on a un p qui s'écrit de deux manières $p = g\pi_1(p_1) = g\pi_2(p_2)$, alors il existe un $p_0 \in \mathbf{U}(R)$ tel que $p_1 = f(p_0)$ et $p_2 = f'(p_0)$. Comme f et f' sont des morphismes, on a $f\sigma_w(p_0) = \tau_w(p_1)$ et $f'\sigma_w(p_0) = \tau'_w(p_2)$ et par commutativité du carré précédent on a $g\pi_1\tau_w(p_1) = g\pi_2\tau'_w(p_2)$. La définition de $\rho_w(p)$ est donc compatible.

γ) Si $p = g\pi_1(p_1)$ on a $\rho_w(p) = g\pi_1\tau_w(p_1)$ de la même forme et $\rho_w^2(p) = g\pi_1\tau_w^2(p_1) = g\pi_1(p_1) = p$. Ainsi, ρ_w est une w-permutation.

δ) On peut reprendre l'essentiel de la construction précédente pour définir ρ_c . Si $p = g\pi_1(p_1)$ avec $p_1 \in P_c(S)$, alors on pose $\rho_c(p) = g\pi_1\tau_c(p_1)$. On procède de même pour S' . Si jamais p s'écrit de deux manières $p = g\pi_1(p_1) = g\pi_2(p_2)$ avec $p_1 \in P_c(S)$ et $p_2 \in P_c(S')$, alors il existe un $p_0 \in \mathbf{U}(R)$ tel que $f(p_0) = p_1$ et $f'(p_0) = p_2$. La propriété de conciliabilité de f et f' s'applique et on déduit d'un coup toute l'orbite de p . On lui associe le symbole commun à toutes les cellules antécédents. Le port principal est alors uniquement déterminée par l'arité.

ϵ) On pose $T = (\rho_w, \rho_c)$ et on a bien $\mathbf{U}(T) = G/\sim$. Par construction on a un carré

commutatif
$$\begin{array}{ccc} & R & \\ f \swarrow & & \searrow f' \\ S & & S' \\ g\pi_1 \searrow & & \swarrow g\pi_2 \\ & T & \end{array}$$
, car T a été construit autour des relations faisant de $g\pi_1$ et $g\pi_2$

des morphismes. Il nous reste à montrer qu'il vérifie la propriété universelle des sommes amal-

gamées. Si on a un autre carré commutatif
$$\begin{array}{ccc} & R & \\ f \swarrow & & \searrow f' \\ S & & S' \\ h \searrow & & \swarrow h' \\ & T' & \end{array}$$
 avec $T = (\rho'_w, \rho'_c)$, en appliquant

$\mathbf{U}$ au diagramme on sait qu'il existe une unique application $G/\sim \xrightarrow{u} \mathbf{U}(T')$ telle que $h = ug\pi_1$ et $h' = ug\pi_2$. Pour conclure il nous faut montrer que u est un morphisme de réseau.

Soit $p \in P_c(T)$, par construction on a, par exemple, $p = g\pi_1(p_1)$ avec $p_1 \in P_c(S)$. On a donc $h(p) \in P_c(T')$ et on en déduit que $u(P_c(T)) \subseteq P_c(T')$

Soit $p \in \mathbf{U}(T)$, supposons que $p = g\pi_1(p_1)$, on a $\rho_w(p) = g\pi_1\tau_w(p_1)$ et donc $u\rho_w(p) = h\tau_w(p_1) = \rho'_w h(p_1) = \rho'_w ug\pi_1(p_1) = \rho_w u(p)$. Là encore, nous n'avons pas utilisé de propriétés

propres aux fils, et le même raisonnement permet de déduire toutes les égalités nécessaires pour être un morphisme. $\blacktriangleright$

2.7.5 Application à la réduction de réseaux

Lemme 2.33 (Contexte) *Si on a*
$$\begin{array}{ccc} & f & R \\ & \swarrow & \\ S & & \\ & \searrow & \\ & g & T \end{array}$$
 alors il existe un S' et une injection $R \xhookrightarrow{f'} S'$

tels que
$$\begin{array}{ccccc} & f & R & & f' \\ & \swarrow & & \searrow & \\ S & & s.a. & & S' \\ & \searrow & & \swarrow & \\ & g & T & & \end{array}$$

Preuve Tout d'abord, on se ramène à des inclusions. En effet, en appliquant deux fois le lemme 2.26 on a le diagramme commutatif suivant :

$$\begin{array}{ccccc} & \tilde{f}(R) & \xrightarrow[\sim]{\tilde{g}} & \tilde{g}\tilde{f}(R) & \\ & \nearrow \tilde{f} & \nearrow \cap & \nearrow \supseteq & \\ R & \xrightarrow[f]{} & S & \xrightarrow[\sim]{\tilde{g}} & \tilde{g}(S) \\ & \searrow g & \searrow \cap & \searrow \supseteq & \\ & & & & T \end{array}$$

Si on arrive en le compléter avec une somme amalgamée à droite :

$$\begin{array}{ccccccc} & \tilde{f}(R) & \xrightarrow[\sim]{\tilde{g}} & \tilde{g}\tilde{f}(R) & & & \\ & \nearrow \tilde{f} & \nearrow \cap & \nearrow \supseteq & \nearrow \supseteq & & \\ R & \xrightarrow[f]{} & S & \xrightarrow[\sim]{\tilde{g}} & \tilde{g}(S) & \text{s.a.} & S' \\ & \searrow g & \searrow \cap & \searrow \supseteq & \searrow \supseteq & & \\ & & & & & & T \end{array}$$

on aura obtenu la somme amalgamée voulue.

Plaçons nous donc dans le cas où $R \subseteq S \subseteq T$. Par définition, on a $S = R \xhookrightarrow{f} \underline{R}$ et $T = S \xhookrightarrow{g} \underline{S}$. On a donc $T = (R \xhookrightarrow{f} \underline{R}) \xhookrightarrow{g} \underline{S}$. En vertu du corollaire 2.8, il existe f_1 et f_2 telles que $T = R \xhookrightarrow{f_1} (\underline{R} \xhookrightarrow{f_2+g} \underline{S})$. Posons $S' = \underline{R} \xhookrightarrow{f_2+g} \underline{S}$.

Pour conclure il nous suffit de montrer qu'il existe un unique morphisme de S' dans la somme amalgamée donnée par le fait 2.28 dans **Set**. En effet, S' devient alors la somme amalgamée au sens ensembliste et la propriété universelle dans **Set** est plus contraignante (il y'a plus de carrés) que dans **IN**.

Or, le morphisme est ici trivial car la relation d'équivalence qui définit la somme amalgamée est en fait l'égalité ⁴. $\blacktriangleright$

⁴Ce qui ne veut pas dire que le quotient est trivial, car c'est la somme disjointe que l'on quotiente.

Définition 2.34 On appelle super-règle la donnée d'un diagramme $R_r \xleftarrow{f_r} R_i \xrightarrow{f_p} R_p$ où R_i n'a pas de cellules et f_r, f_p sont superposables.

Théorème 2.35 Si $R_r \xleftarrow{f_r} R_i \xrightarrow{f_p} S_p$ est une super-règle et qu'on a un morphisme $R_r \xrightarrow{g} R$ alors on peut faire la complétion de diagramme suivante

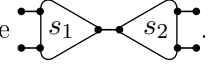
$$\begin{array}{ccccc} R_r & \xleftarrow{f_r} & R_i & \xrightarrow{f_p} & R_p \\ g \downarrow & & \downarrow & & \downarrow \\ R & \xleftarrow{\quad} & S & \xrightarrow{\quad} & T \end{array}$$

s.a. *s.a.*

T est alors appelé réduit de R par la super-règle.

Preuve Tout d'abord, on peut appliquer le lemme 2.33 pour obtenir la somme amalgamée de gauche. Par construction, f_r et $R_i \xrightarrow{g} S$ sont superposables. Comme f_r et f_p sont fortement non superposables, on en déduit que f_p et g sont superposables, a fortiori conciliables. On peut alors appliquer le théorème 2.32 qui nous donne la somme amalgame de droite.

La réduction ainsi définie coïncide avec la précédente lorsque R_r est de la forme



Chapitre 3

Boîtes

Les boîtes sont antérieures aux réseaux d'interaction, elles furent introduites pour rendre compte de points impossibles à déséquentialiser dans les réseaux de preuve.

Il est immédiat d'adapter leur définition au cadre plus général des réseaux d'interaction, et leur utilisation y est aussi très naturelle. En effet, imaginons que nous voulons faire passer à travers toute une chaîne de réduction un réseau comme argument, il serait pratique d'*emballer* ce réseau dans une cellule, aisément passable en argument ou duplicable. De plus, une telle construction permettrait de garantir une certaine sûreté de la réduction dans la mesure où aucune interaction ne serait possible avant l'application d'une opération de *déballage*.

Notons qu'une tradition similaire existe dans les langages de programmation usuels : les types *emballés* (en anglais *boxed types*). Le but est alors le même : considérer des données de manière complètement opaque.

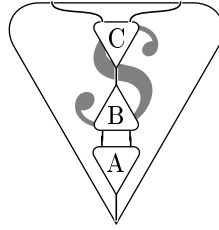
3.1 Définition

Nous notons $\mathcal{R}_0$ l'ensemble des réseaux d'interaction. Nous allons définir maintenant l'ensemble des réseaux d'interaction à boîtes, noté $\mathcal{R}$, comme le plus petit point fixe d'une construction.

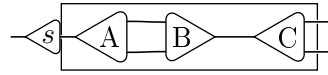
Soit E un ensemble de réseaux d'interaction (à boîtes), on appelle réseau d'interaction à boîtes dans E un couple (R, β) où R est un réseau d'interaction et β est un étiquetage partiel des cellules, tel que pour $c \in \text{dom}(\beta)$, $\beta(c) = (r, i)$ où $r \in E$ et i est une bijection de $P_f(r)$ dans c . La cellule c est appelée une *boîte*. La condition d'arité force toutes les boîtes d'un symbole donné à ne contenir que des réseaux ayant un nombre de ports libres fixé. Cette restriction est aisément contournée en considérant une famille de symboles $(s_i)_{i \in \mathbb{N}}$ comme véritable symbole de la boîte. Par abus de langage, nous parlerons juste de symbole pour faire référence à cette famille.

Notons f l'application passant de E à l'ensemble des réseaux à boîtes dans E . On pose $\mathcal{R} = \biguplus_i \mathcal{R}_i$ où les $\mathcal{R}_i$ sont construits à partir de $\mathcal{R}_0$ et $\mathcal{R}_{i+1} = f(\mathcal{R}_i)$.

On pourra représenter une boîte comme une cellule dans laquelle on dessine le réseau étiquette et où la bijection i est directement visible :



Cependant, ce type de représentation n'est pas très lisible et on lui préférera le type suivant, où l'on encadre le réseau contenu par une boîte. Une cellule représente alors le port principal et les ports auxiliaires sont ordonnés de manière quelconque. Dans la pratique, lorsque l'ordre sera important on étiquettera ces ports auxiliaires, levant ainsi l'ambiguïté.



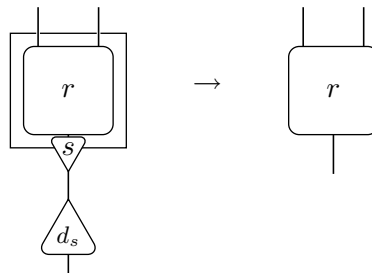
Le port libre correspondant au port principal (resp. à un port auxiliaire) de la boîte sera appelé *porte principale* (resp. *porte auxiliaire*).

3.2 Opérations de base sur les boîtes

Dans ce paragraphe, nous considérons un symbole de boîte, noté s , et nous allons définir des opérations naturelles sur les boîtes correspondantes.

3.2.1 Déballage

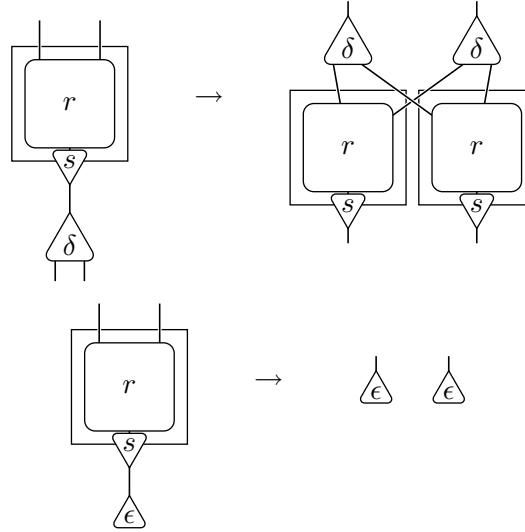
L'opération essentielle pour manipuler des boîtes consiste à déballer le contenu de la boîte dans le réseau qui la contient. On note d_s le symbole de la cellule de déballage, avec $\alpha(d_s) = 1$. La règle de réduction correspondante a la forme suivante.



Notons que cela n'est pas une règle de réduction au sens usuel. Le motif à droite de la règle ne dépend pas juste du couple (s, d_s) mais également du contenu de la boîte.

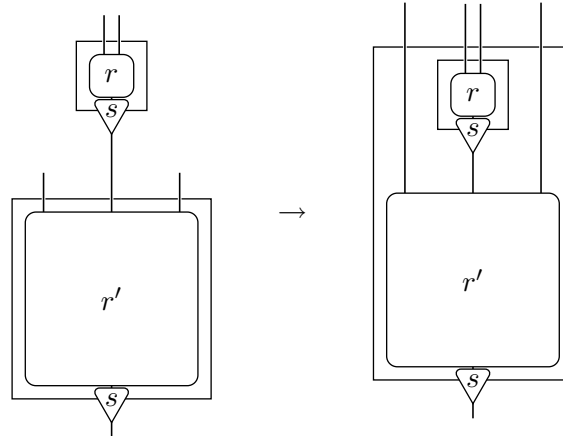
3.2.2 Duplication et effacement

On va maintenant définir deux opérations, la duplication, notée δ , et l'effacement, noté ϵ , que l'on va considérer intuitivement comme génériques : c'est-à-dire définie par réduction contre toute autre cellule. En effet, on souhaiterait en effaçant une boîte, effacer également tout ce qui y est connecté. On a donc les deux réductions suivantes.



3.2.3 Enfouissement

On appelle *enfouissement* (de l'anglais *digging*) la réduction suivante.



L'intérêt d'une telle règle n'est pas immédiat, et ne sera justifié que plus tard lorsque nous nous intéresserons à la logique linéaire. Cependant, nous pouvons, d'ores et déjà remarquer que cette règle sort du cadre des règles d'interaction. Cette modification a l'air bénigne : par rapport à la définition usuelle du motif d'une règle d'interaction, on a juste remplacé un fil *port principal sur port principal* par un fil *port principal sur port auxiliaire*. Pourtant, cela suffit à ne plus assurer la confluence forte.

Il est toutefois possible de montrer la confluence de la réduction. Les figures 3.1 et 3.2 présentent deux des paires critiques, les deux autres étant aussi résolubles de manière immédiate.

3.2.4 Agrégation

La présence des opérations de déballage et d'enfouissement donne une structure très complexe à la manipulation des boîtes. Une approche plus simple est de fusionner ces deux opérations : une boîte ne s'enfouit plus dans une autre, à la place les deux boîtes mettent en

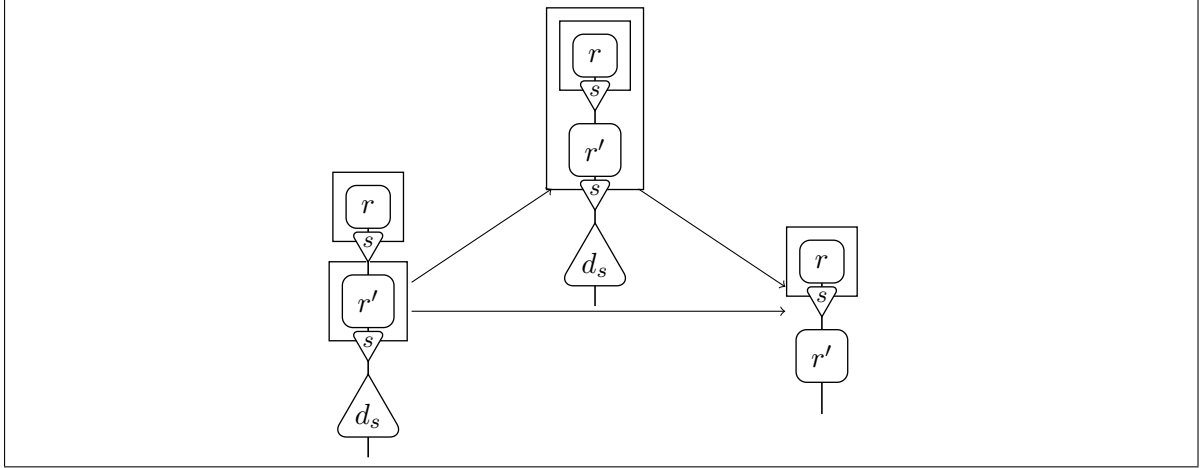


Figure 3.1: Résolution de la paire critique *débailage/enfouissement*

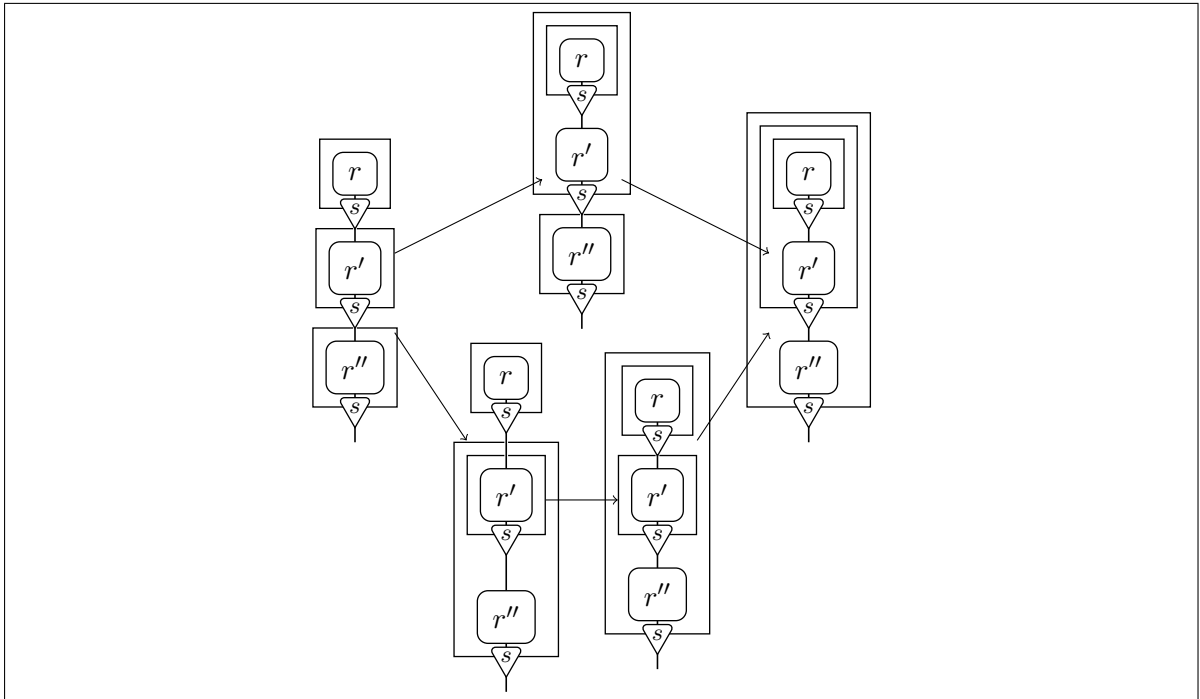
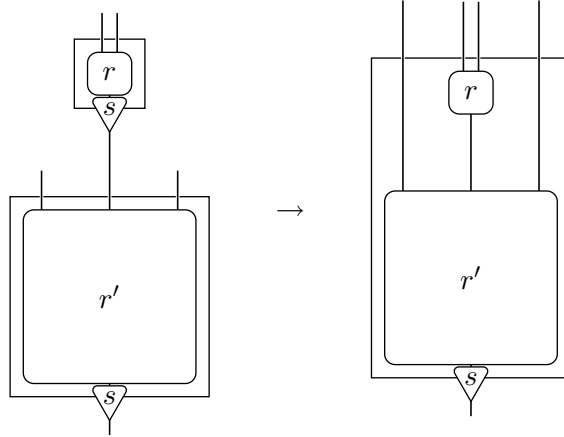


Figure 3.2: Résolution de la paire critique *enfouissement/enfouissement*

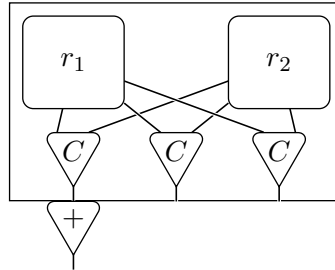
commun leur contenu.



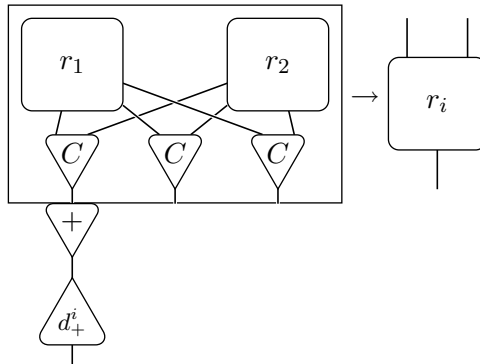
Remarque 3.1 Avec cette opération également, la réduction n'est plus assurée d'être fortement confluente. Cependant, on peut prouver la confluence vis-à-vis de la duplication et de l'effacement.

3.3 Boîtes additives

On introduit maintenant un raffinement sur les boîtes qui va nous permettre de contenir la superposition de deux réseaux. On se fixe un symbole binaire C qui va permettre de *contracter* les ports des deux réseaux. On appellera boîte additive une boîte de la forme



Les boîtes additives peuvent être vues comme une simple restriction sur leur contenu en tant que boîte. Cette restriction permet d'énoncer des règles ne déballant qu'un des deux réseaux :



Chapitre 4

Réseaux d'interaction et logique linéaire

Dans ce chapitre nous allons présenter des réseaux de preuve pour différents fragments de la logique linéaire. Comme notre étude se place dans le cadre des réseaux d'interaction, nous n'aurons pas exactement des réseaux de preuve mais leur Ex-effondrement (cf. paragraphe 2.6).

4.1 Bibliothèques closes

Soit $\mathcal{S}' \subseteq \mathcal{S}$ un sous-ensemble de symboles, notons $\mathfrak{R}(\mathcal{S}')$ l'ensemble des réseaux dont les cellules ont leur symboles dans $\mathcal{S}'$.

On rappelle qu'une bibliothèque est un ensemble de règles d'interaction ne comportant qu'au plus une règle par paire de symboles (cf. corollaire 2.18), rappelons également que la réduction y est alors fortement confluente.

Définition 4.1 *Soit $\mathcal{S}' \subseteq \mathcal{S}$ et $\mathcal{B}$ une bibliothèque.*

$\mathcal{B}$ est dite close pour $\mathcal{S}'$ lorsque pour toute règle $\mathcal{R} = (R_r^{I_r}, R_p^{I_p}) \in \mathcal{B}$ on a $R_r, R_p \in \mathfrak{R}(\mathcal{S}')$.

Fait 4.2 *Si $\mathcal{B}$ est une bibliothèque close pour $\mathcal{S}'$, alors $\rightarrow_{\mathcal{B}}$ est une réduction sur $\mathfrak{R}(\mathcal{S}')$.*

La notion de bibliothèque close nous permet donc de capturer la notion de fragment logique ou de sous-modèle de calcul.

4.1.1 Dans le cadre des réseaux d'interactions avec boîtes

Soit $\mathcal{S}', \mathcal{S}_b \subseteq \mathcal{S}$, notons $\mathfrak{R}(\mathcal{S}', \mathcal{S}_b)$ l'ensemble des réseaux dont les cellules sont étiquetées par $\mathcal{S}'$ et les boîtes par $\mathcal{S}_b$. On supposera toujours $\mathcal{S}' \cap \mathcal{S}_b = \emptyset$.

On adapte à l'identique la définition précédente en substituant $\mathfrak{R}(\mathcal{S}', \mathcal{S}_b)$ à $\mathfrak{R}(\mathcal{S}')$, ce qui nous permet de définir la notion de bibliothèque close pour des réseaux d'interactions avec boîtes.

4.2 Typage de réseaux

Nous allons ajouter de l'information aux réseaux en typant les ports par des formules logiques. Nous considérons donc ici un ensemble de formule $\mathcal{F}$ que l'on considère clos par une

involution $\square^\perp$.

Un fil $A \multimap B$ est dit bien typé lorsque $B = A^\perp$. La donnée des deux types étant alors redondante on notera juste $A \multimap$ où la flèche pointe vers le port du type indiqué.

On type les cellules avec des règles de typage de la forme



Bien que le sens de telles règles soit en général sans ambiguïté, on peut donner une définition plus formelle.

Étant donné $\mathcal{S}' \subseteq \mathcal{S}$, pour chaque symbole $s \in \mathcal{S}'$ on se donne un couple (l, f) où l est un $\alpha(s)$ -uplet $(\mathcal{F}_1, \dots, \mathcal{F}_{\alpha(s)})$ de sous-ensembles de $\mathcal{F}$, représentant les types possibles des ports auxiliaires associés, et $f : l \rightarrow \mathcal{F}$ donne le type du port principal en fonction du type des ports auxiliaires.

On dira qu'un réseau est *bien typé* vis-à-vis d'un ensemble de règles de typage lorsque toutes ces cellules respectent les règles de typage et que tous les fils sont bien typés.

Soit $\mathcal{B}$ une bibliothèque close pour $\mathcal{S}'$. Le typage nous permet de caractériser les formes normales de $\rightarrow_{\mathcal{B}}$ sous réserve d'une certaine complétude.

Définition 4.3 Une bibliothèque $\mathcal{B}$ close pour $\mathcal{S}'$ est dite complète vis-à-vis du typage si tout fil actif d'un réseau de $\mathfrak{R}(\mathcal{S}')$ bien typé peut être réduit par $\mathcal{B}$.

En effet, on a le résultat direct suivant :

Propriété 4.4 Soit $\mathcal{B}$ complète vis-à-vis du typage, si $R \in \mathfrak{R}(\mathcal{S}')$ bien typé est en forme normale alors il est en forme normale absolue.

En d'autres termes, le typage nous assure que la réduction n'a pas de blocages, de *deadlocks*.

4.3 MLL

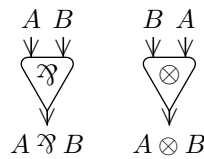
Le fragment le plus simple de la logique linéaire est la *logique linéaire multiplicative*, ou MLL.

Les formules de MLL sont données par la grammaire suivante

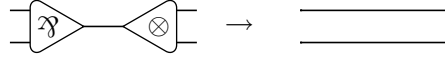
$$\mathcal{F} := V_P | V_P^\perp | \mathcal{F} \wp \mathcal{F} | \mathcal{F} \otimes \mathcal{F}$$

où V_P est un ensemble dénombrable de variables propositionnelles. On étend syntaxiquement l'involution $^\perp$ à tout $\mathcal{F}$ par $(A \wp B)^\perp = A^\perp \otimes B^\perp$.

Ce fragment a pour connecteurs $\wp$ et $\otimes$, avec $\alpha(\wp) = \alpha(\otimes) = 2$ et

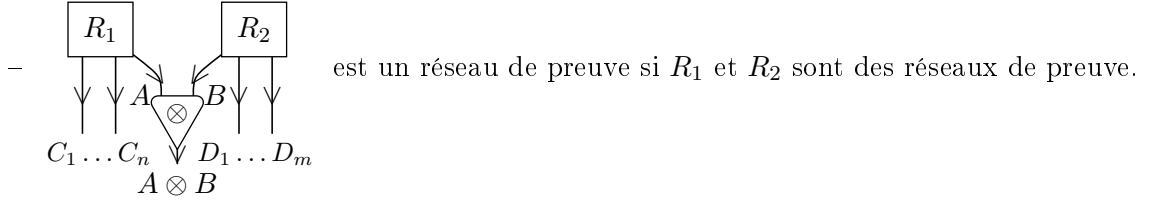


Et on a l'unique règle de réduction suivante :

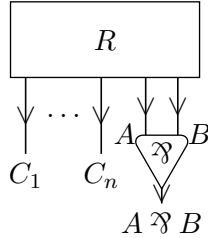


Parmi ces réseaux d'interaction bien typé on distingue les réseaux de preuve comme étant définis de manière inductive par :

- $A \Leftarrow$ est un réseau de preuve



- R est un réseau de preuve si R est un réseau de preuve.



Il est possible de caractériser les réseaux de preuve parmi les réseaux d'interaction de MLL de manière géométrique. Soit $R = (\sigma_w, \sigma_c)$ un réseau de MLL, on appelle *switching* tout graphe non orienté (V, E) tel que

- $V = P(R)$
- $\forall p \in P(R), \{p, \sigma_w(p)\} \in E$
- si c est une cellule de symbole $\otimes$, $\forall i \in \{1, 2\}, \{\text{pal}(c), \text{pax}_i(c)\} \in E$
- si c est une cellule de symbole $\wp$, on choisit $i \in \{1, 2\}$ et $\{\text{pal}(c), \text{pax}_i(c)\} \in E$

Si le réseau comporte n cellules $\wp$ il admet donc 2^n switchings différents.

On a le théorème suivant dû à Danos et Regnier, dont une preuve pourra être trouvée dans [DR89].

Théorème 4.5 (Critère de Danos-Regnier) *Soit R un réseau de MLL. R est un réseau de preuve si et seulement si tous ses switchings sont acycliques connexes.*

4.4 MELL

MLL est assez limité : toutes les ressources manipulées sont linéaires. La *logique linéaire multiplicative exponentielle*, MELL, consiste à ajouter à MLL un système de boîtes : les !-boîtes, encore appelées *promotion*, auxquels on ajoute les opérations de base décrites au paragraphe 3.2.

La duplication sera appelée *contraction*, le déballage : *déréliction* et l'effacement : *affaiblissement*. De plus on note abusivement ? le symbole de ces trois opérations. Dans la mesure où elles sont d'arité distinctes il n'y a pas de confusion possible. Cependant, lorsque on voudra faire explicitement référence à l'un de ces symboles, on notera $?^0$ le symbole d'affaiblissement, $?^1$ celui de la déréluction et $?^2$ celui de la contraction.

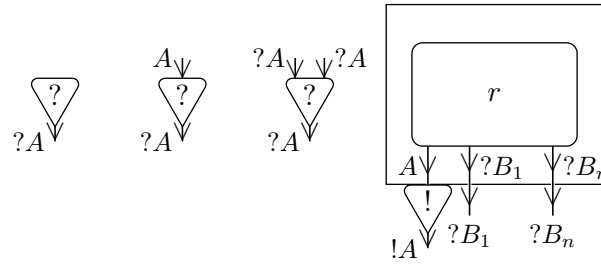
Du fait de la présence de la règle d'enfouissement, la réduction de MELL est confluente mais pas fortement confluente.

Les formules de MELL sont données par la grammaire suivante, étendant les formules de MLL :

$$\mathcal{F} := V_P | V_P^\perp | \mathcal{F} \wp \mathcal{F} | \mathcal{F} \otimes \mathcal{F} | !\mathcal{F} | ?\mathcal{F}$$

où V_P est un ensemble dénombrable de variables propositionnelles. On étend syntaxiquement l'involution $^\perp$ à tout $\mathcal{F}$ par $(A \wp B)^\perp = A^\perp \otimes B^\perp$ et $(!A)^\perp = ?A^\perp$.

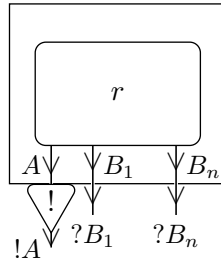
Aux règles de typage de MLL, on ajoute :



4.4.1 ELL

Le fragment ELL, pour *Elementary Linear Logic*, correspond à MELL où les règles d'enfouissement et de déballage se sont vues remplacées par la règle d'agrégation. Bien que nous ne préoccupons pas ici de complexité, notons que ELL a été introduit afin de caractériser la classe des fonctions calculable en temps élémentaires : toute fonction calculable en temps élémentaire est représentable par une preuve de ELL dont le calcul s'effectue au cours de la réduction, et de plus la longueur d'une telle réduction est au plus élémentaire.

La règle de typage de la promotion doit être changée pour que le typage reste préservé par réduction, on la remplace par :



4.4.2 Connecteurs synthétiques

Une syntaxe plus synthétique pour les réseaux de MELL est définie dans la thèse de Laurent Regnier [Reg92]. Celle-ci part du principe que l'on peut faire passer toutes les contractions en dehors des boîtes et elle consiste à remplacer un arbre de contractions dont les n feuilles sont des déréllections ou des affaiblissements par une unique cellule n -aire. Deux tels arbres ayant un même nombre de feuilles sont remplacés par la même cellule. Ainsi, cette syntaxe est un quotient de la syntaxe d'origine. On appellera réseaux de SMELL (*Synthetic MELL*) de tels réseaux.

Notons que les déréllections sont maintenant superflues et le nouveau connecteur n -aire, appelé *arbre exponentiel*, a le typage

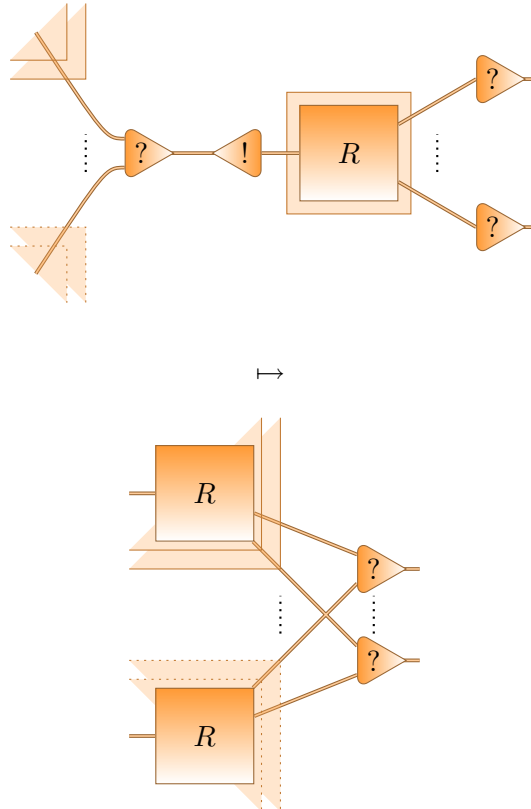


De plus, on demande à ce que les portes auxiliaires de boîtes soient reliées à des ports auxiliaires d'arbres exponentiels.

On appelle *hauteur* d'une cellule le nombre de boîtes contenant celle-ci. On appelle *hauteur* d'un port auxiliaire d'un arbre exponentiel le nombre de boîtes contenant un fil atteint depuis ce port soit en progressant le long des fils soit en rentrant dans les boîtes. On appelle *portance* d'un port auxiliaire d'un arbre exponentiel la différence entre sa hauteur et la hauteur de la cellule.

Remarque 4.6 *La définition de la hauteur d'un port auxiliaire semble inutilement compliquée, en fait la situation est telle du fait de l'absence de liens axiome. En effet, dans les réseaux de preuve il suffit de dire que la hauteur d'un port auxiliaire est la hauteur du lien directement accessible depuis ce port.*

La règle de réduction portant sur les arbres exponentiels est naturellement définie comme celle obtenue en le remplaçant par un arbre de contractions dont les feuilles sont reliées à des déréllections placées à la profondeur de boîte associé à la hauteur :



Deuxième partie

Géométrie de l'Interaction

Chapitre 5

Semigroupes inversifs

Dans ce chapitre, on suivra la tradition anglophone, utilisant *semigroupe* pour le concept anglais de *semigroup* et le concept français de *demi-groupe*. Notre raison principal en est que les logiciens ont l'habitude de les dénommer ainsi pour parler de géométrie de l'interaction. De plus, étant donné que nous ne ferons jamais référence à la notion française de semi-groupe ¹ il n'y aura pas d'ambiguïté possible.

La plupart des résultats présentés ici sont classiques, on en trouvera une exposition complète dans un livre de référence comme [How95].

Le but de ce chapitre est double. D'un coté on donne tous les outils nécessaires à la future définition de la Géométrie de l'Interaction et à son étude : théorie équationnelle, modèle, ..., et d'un autre coté on montre un lien très fort entre ces approches et l'approche originelle de Girard à l'aide d'algèbre stellaire. En effet, à la fin du chapitre on prouve que l'algèbre associée à un semigroupe inversif est représentable en tant qu'algèbre d'opérateurs sur un espace de Hilbert, et que cette représentation coïncide avec celle donnée par Girard. Cela justifie donc complètement notre focalisation sur les semigroupes inversifs : ils représentent le noyau mathématique de la GdI.

5.1 Définitions

5.1.1 Axiomatique

Définition 5.1 *On appelle semigroupe un ensemble D muni d'une loi de composition $\circ$ associative, c'est-à-dire telle que*

$$a \circ (b \circ c) = (a \circ b) \circ c \quad \forall a, b, c \in D$$

Dans la suite on notera $\circ$ implicitement.

Définition 5.2 *Soit D un semigroupe, on appelle*

- neutre à droite, ou unité à droite, un élément $e \in D$ tel que $xe = x$, $\forall x \in D$
- zéro à droite, un élément $z \in D$, tel que $xz = z$, $\forall x \in D$
- idempotent, un élément $i \in D$, tel que $i^2 = i$
- inverse d'un élément $x \in D$, un élément $a \in D$ tel que $a = axa$ et $x = xax$

¹monoïde régulier

On définit dualement les notions à gauche. Un neutre à gauche et à droite, sera appelé simplement neutre, de même pour les autres propriétés.

On note $\mathfrak{E}(\mathbf{D})$ l'ensemble des idempotents contenus dans $\mathbf{D}$.

Définition 5.3 Un semigroupe est appelé monoïde quand il possède une unité. On la note usuellement 1.

Un semigroupe est dit inversif quand chacun de ses éléments possède une unique inverse, qu'on note a^* pour a .

Un monoïde étant également inversif en tant que semigroupe est appelé monoïde inversif.

Un monoïde inversif où l'inverse satisfait de plus $aa^* = a^*a = 1$, pour tout a , est appelé groupe.

La définition donnée ici de semigroupe inversif par existence et unicité de l'inverse n'est pas la définition *standard* que l'on peut en trouver. En effet, la définition usuelle est celle d'un magma associatif muni d'une application $\cdot^*$ telle que

$$(a^*)^* = a \quad (ab)^* = b^*a^* \quad aa^*bb^* = bb^*aa^* \quad a^*aa^* = a$$

On va voir que notre définition implique ces relations, en fait, on a également la réciproque.

Fait 5.4 Les idempotents d'un semigroupe inversif commutent.

Preuve Soient u et v deux idempotents de $\mathbf{D}$ un semigroupe inversif, on pose $x = (uv)^*$, on a

$$(uv)(vXu)(uv) = uv \quad (vXu)(uv)(vXu) = vXu$$

Donc, vXu est inverse de uv , l'inverse étant unique on a $x = vXu$ et donc x est idempotent car $(vXu)^2 = vXuvXu = vXu$. Il donc son propre inverse et on a $x = (uv)^* = uv$ par unicité. Un argument similaire donne $vu = (vu)^*$. En remarquant que

$$(uv)(vu)(uv) = (uv)^2 = uv \text{ et } (vu)(uv)(vu) = (vu)^2 = vu$$

on a donc $uv = (uv)^* = vu$ par unicité. ◀

Lemme 5.5 Soit $\mathbf{D}$ un semigroupe inversif, pour tout $a \in \mathbf{D}$, aa^* et a^*a sont idempotents.

Preuve Soit $a \in \mathbf{D}$, on a $a(a^*aa^*) = aa^*$ par définition de l'inverse, et donc aa^* est idempotent. De même, on a $(a^*aa^*)a = a^*a$. ▶

Fait 5.6 Soit $\mathbf{D}$ un semigroupe inversif, on a

$$i) \quad \forall a \in \mathbf{D}, (a^*)^* = a$$

$$ii) \quad \forall a, b \in \mathbf{D}, (ab)^* = b^*a^*$$

Preuve

i) Soit $a \in \mathbf{D}$, on a $a^*aa^* = a$ et $aa^*a = a^*$, donc a est un inverse de a^* . Par unicité de l'inverse on en déduit $(a^*)^* = a$.

ii) Soient $a, b \in D$, on a

$$(ab)(b^*a^*)(ab) = a(bb^*)(a^*a)b$$

mais bb^* et a^*a étant idempotents (lemme 5.5), ils commutent, et

$$(ab)(b^*a^*)(ab) = aa^*abb^*b = ab$$

$$(b^*a^*)(ab)(b^*a^*) = b^*(a^*a)(bb^*)a^* = b^*bb^*a^*aa^* = b^*a^*$$

On a donc $(ab)^* = b^*a^*$ par unicité.



5.1.2 Exemples canoniques

Dans ce paragraphe on fixe un ensemble X et on s'intéresse à la structure de certains ensembles d'applications de X dans lui-même.

On pose :

- $\mathcal{T}_X$ l'ensemble des applications totales de X dans X ,
- $\mathcal{P}_X \supset \mathcal{T}_X$ l'ensemble des applications partielles de X dans X , on note alors, pour f dans $\mathcal{P}_X$, $\text{dom}(f)$ le domaine de f et $\text{codom}(f)$ son image,
- $\mathcal{IP}_X \subset \mathcal{P}_X$ l'ensemble des injections partielles de X dans X ,
- $\mathfrak{S}_X \subset \mathcal{T}_X$ l'ensemble des applications bijectives de X dans X .

On considère tous ces ensembles comme des semigroupe muni de la loi associative $\circ$.

Fait 5.7 – $\mathcal{P}_X$ est un semigroupe

- $\mathcal{T}_X$ est un sous-semigroupe de $\mathcal{P}_X$
- $\mathfrak{S}_X$ est un groupe
- $\mathcal{IP}_X$ est un monoïde inversif

5.1.3 Présentation par générateurs et relations

On considère ici connue la construction d'un semigroupe par générateurs et relations. Notons $[g_1, \dots, g_n \mid r_1, \dots, r_m]$ le semigroupe engendré par les g_i et les relations r_j . Pour définir un semigroupe inversif ainsi on peut se ramener à la construction d'un semigroupe pour autant que l'on rajoute des nouveaux générateurs étoilés et les relations adéquats. Remarquons qu'il s'agit de la même technique utilisée pour définir un groupe par présentation en tant que monoïde.

Ainsi on pose :

$$\langle g_1, \dots, g_n \mid r_1, \dots, r_m \rangle = [g_1, g_1^*, \dots, g_n, g_n^* \mid r_1, \dots, r_m, r_1', r_1'', r_1''', \dots, r_n', r_n'', r_n''']$$

où $r_i' ::= g_i = g_i g_i^* g_i$, $r_i'' ::= g_i^* = g_i^* g_i g_i^*$ et $r_i''' ::= (g_i^*)^* = g_i$.

Définition 5.8 Soit $D = \langle \bar{g} \mid \bar{r} \rangle$ un semigroupe inversif (avec zéro) présenté par générateurs et relations, on appelle semigroupe positif de D le semigroupe engendré par les g_i ($\neq 0$), on le note D^+ .

Un élément positif est donc un élément qui admet une écriture sans $\star$.

On note $\langle \bar{g} \mid \bar{r} \rangle_0 = \langle \bar{g}, 0 \mid \bar{r}, g_1 0 = 0, 0 g_1 = 0, \dots, g_n 0 = 0, 0 g_n = 0 \rangle$, ce qui permet d'exprimer directement la présentation de semigroupe inversif avec zéro. De même, on note $\langle \bar{g} \mid \bar{r} \rangle_0^1$ la présentation d'un monoïde inversif avec zéro.

Définition 5.9 Soit $D = \langle \bar{g} \mid \bar{r} \rangle_0$, on dit que D est normal si $\forall a, b \in D^+ \quad ab^* \neq 0$

La réciproque de cette propriété est traditionnellement appelée *propriété ab^** :

Définition 5.10 Soit $D = \langle \bar{g} \mid \bar{r} \rangle$ un semigroupe inversif normal, on dit que D a la propriété ab^* si

$$\forall x \in D, x \neq 0 \Rightarrow \exists a, b \in D^+ \quad x = ab^*$$

On substituera dans la plupart des cas une notion plus forte à la normalité :

Définition 5.11 Soit $M = \langle \bar{g} \mid \bar{r} \rangle_0^1$ un monoïde inversif avec zéro, on dit que M est simplifiable lorsque $\forall a \in M^+, a^*a = 1$.

Fait 5.12 Soit $M = \langle \bar{g} \mid \bar{r} \rangle_0^1$ un monoïde inversif avec zéro où $0 \neq 1$, alors M simplifiable entraîne M normal.

Preuve Soient $a, b \in M^+$, supposons $ab^* = 0$, alors $a^*ab^*b = 0$ mais M étant simplifiable on en déduit $1 = 0$, ce qui est contradictoire. $\blacktriangleleft$

5.1.3.a Automorphismes

Il peut être pratique de considérer en plus des générateurs un certain nombre d'automorphismes dans une présentation. Par exemple on pourra noter

$$\langle x \mid f \mid \rangle$$

le semigroupe librement engendré par x et stable par l'application de f .

Définition 5.13 On note $D = \langle \bar{g} \mid \bar{f} \mid \bar{r} \rangle$ le semigroupe $\bigcup_i E_i / R$ où $E_0 = \langle \bar{g} \rangle$ et

$$E_{i+1} = \bigcup_k \{f_k(x) \mid x \in E_i\}$$

muni du produit $f_i(x)f_i(y) = f_i(xy)$ déduit du produit sur E_0 , et où R est la relation engendrée par les $\bar{r}$.

Notons que dans cette définition les relations peuvent faire intervenir les automorphismes, c'est d'ailleurs son principal intérêt.

Dans le cas d'un semigroupe inversif, on note également D^+ le semigroupe des éléments positifs de D , et on rajoute la relation $f_i(x^*) = f_i(x)^*$. De même, lorsqu'ils sont définis on pose $f_i(0) = 0$ et $f_i(1) = 1$.

5.1.3.b Modèles

Il est souvent plus intuitif de considérer une présentation par générateurs et relations comme une théorie équationnelle dont on cherche des modèles. Le semigroupe effectivement engendré par cette présentation étant le modèle le plus général possible satisfaisant la théorie équationnelle. Nous avons choisi ici de ne pas partir de cette notion de théorie équationnelle, cependant il est possible de récupérer la notion de modèle.

Définition 5.14 Soit $M = \langle \bar{g} \mid \bar{r} \rangle_0^1$, on appelle modèle de M un monoïde inversif avec zéro M' muni d'un morphisme de monoïde avec zéro $\phi : M \rightarrow M'$. On note $M' \models M$ et si $\phi(x) = \phi(y)$ on note $M' \models x = y$.

Un modèle M' de M est dit

- non-trivial lorsque $M' \not\models 0 = 1$
- fidèle lorsque $M' \models x = 0$ si et seulement si $x = 0$

Le principal intérêt des modèles provient des deux faits suivants :

Fait 5.15 Si $M = \langle \bar{g} \mid \bar{r} \rangle_0^1$ admet un modèle non-trivial alors $0 \neq 1$ dans M .

Preuve Trivialement, si on avait $0 = 1$ alors pour tout modèle M' on aurait $M' \models 0 = 1$. $\blacktriangleright$

Fait 5.16 Si $M = \langle \bar{g} \mid \bar{r} \rangle_0^1$ a la propriété ab^* et qu'il est simplifiable alors tout modèle non-trivial est fidèle.

Preuve Soient M' un modèle non-trivial de M et $x \in M$ tel que $M' \models x = 0$. Si $x \neq 0$ alors $x = ab^*$ et $1 = a^*ab^*b$, qui donne $M' \models 1 = a^*xb = 0$ contradictoire. On a donc $x = 0$. $\blacktriangleright$

5.2 Propriétés essentielles

5.2.1 Ordre

Étant donné un semigroupe inversif D on peut y définir un ordre partiel $\leq$ par

$$a \leq b \iff \exists e \in \mathfrak{E}(D), a = eb$$

Fait 5.17 La relation $\leq$ est un ordre partiel sur D .

Preuve

- i) Soit $a \in D$, $a = aa^*a$ et $aa^* \in \mathfrak{E}(D)$ donc $a \leq a$
- ii) Soient $a, b \in D$ tels que $a \leq b$ et $b \leq a$, i.e. $a = eb$ et $b = fa$ avec $e, f \in \mathfrak{E}(D)$. On a

$$a = eb = efa = fea = feeb = feb = fa = b$$

- iii) Soient $a, b, c \in D$ tels que $a \leq b$ et $b \leq c$, i.e. $a = eb$ et $b = fc$ avec $e, f \in \mathfrak{E}(D)$. On a donc $a = (ef)c$ mais $ef \in \mathfrak{E}(D)$, donc $a \leq c$. $\blacktriangleright$

5.2.2 Structure des idempotents

Les idempotents d'un semigroupe inversif D ont une structure remarquable que l'on va étudier ici.

Le lemme suivant, essentiel dans la future démonstration du théorème de Preston-Wagner, présente quelques résultats sur la structure des classes à droite induites par des idempotents.

Lemme 5.18 Soit D un semigroupe inversif contenu dans un semigroupe S .

- i) Soient $e, f \in \mathfrak{E}(D)$, $Se = Sf \Rightarrow e = f$.
- ii) Soient $e, f \in \mathfrak{E}(D)$, $Se \cap Sf = Se f$.

iii) Soient $x \in D$, $Sx^*x = Sx$.

Les propriétés symétriques définies par multiplication à gauche sont aussi vraies, les preuves étant symétriques.

Preuve (i) $e, f \in D \subseteq S$ donc $e = ee = xf$ avec $x \in S$. $ef = xff = xf = e$. De même, on a $fe = f$, mais $ef = fe$ par 5.4, et donc $e = f$.

(ii) L'inclusion $Se f \subseteq Se \cap Sf$ découle directement de 5.4. Soient, $xe = yf \in Se \cap Sf$, on a $xef = xef = yff = yf = xe$, donc $xe \in Se f$ et on conclut l'égalité.

(iii) On a $Sx^*x \subseteq Sx = Sxx^*x \subseteq Sx^*x$. ◀

Fait 5.19 $\mathfrak{E}(D)$ forme un demi-treillis inférieur.

Par demi-treillis inférieur, on entend que pour toute pair $(e, f) \in \mathfrak{E}(D)^2$, il existe un élément $e \wedge f \in \mathfrak{E}(D)$, appelé borne inférieure, satisfaisant

$$\forall g \in \mathfrak{E}(D), g \leq e \text{ et } g \leq f \Rightarrow g \leq e \wedge f$$

Preuve Soient $e, f \in \mathfrak{E}(D)$, montrons que $e \wedge f = ef$ convient. Notons déjà que $e \wedge f = f \wedge e$.

Soit $z \in \mathfrak{E}(D)$ tel que $z \leq e$ et $z \leq f$, cela signifie que $z = e'e = f'f$, on a donc $z = e'e = e'ee = ze = e f'f$ et $z \leq e \wedge f$. ◀

5.2.3 Théorèmes de plongement

On appelle théorème de plongement pour une structure donnée, un théorème qui fournit un plongement d'un objet de cette structure vers une sous-structure d'un objet canonique. L'exemple le plus connu d'un tel théorème est celui de Cayley, plongeant un groupe dans le groupe de ses permutations. Le fait qu'un tel théorème soit également démontrable dans le cas des semigroupes représente un des intérêts principaux pour cette structure.

Dans cette section nous allons énoncer et démontrer ces deux théorèmes. L'image par le plongement du théorème de Cayley sera appelé par la suite représentation de Cayley, de même on définira la représentation de Preston-Wagner.

Théorème 5.20 (Cayley) Soit G un groupe, il existe un morphisme injectif de groupe de G dans $\mathfrak{S}_G$.

Preuve Soit $g \in G$, on note ρ_g l'application défini par

$$\rho_g(g') = gg'$$

On va montrer que ρ_g est une bijection, et donc $\rho_g \in \mathfrak{S}_G$. Puis pour conclure, que ρ est un morphisme injectif.

i) L'injectivité de ρ_g est directe :

$$gg' = gg'' \Rightarrow (g^{-1}g)g' = (g^{-1}g)g'' \Rightarrow g' = g''$$

Soit $g' \in G$, on a $g' = g(g^{-1}g') = \rho_g(g^{-1}g')$. On a donc $\rho_g \in \mathfrak{S}_G$.

ii) Soient $g, g', g'' \in G$, on a $(\rho_g \circ \rho_{g'})(g'') = gg'g'' = \rho_{gg'}(g'')$ et $\rho_1(g'') = 1g'' = g''$. ρ est un morphisme de groupe.

Supposons maintenant que $\rho_g = \rho_{g'}$ pour $g, g' \in G$. Cela signifie en particulier $\rho_g(1) = \rho_{g'}(1)$, i.e. $g = g'$. ρ est donc un morphisme injectif.²

²Une autre manière de le voir est de dire qu'un élément du noyau de ρ est nécessairement un neutre, et par unicité ce noyau est trivial.



Théorème 5.21 (Preston-Wagner) *Soit D un semigroupe inversif, il existe un morphisme injectif de semigroupe de D dans $\mathcal{IP}_D$.*

Preuve La preuve de ce théorème est très proche de celle du théorème de Cayley, mis à part une difficulté accrue du fait de la structure plus lâche de semigroupe inversif et d'injection partielle.

Soit $x \in D$, on note φ_x l'application partielle de domaine $x^*D = \{x^*y \mid y \in D\}$ définie par $\varphi_x(y) = xy$. On a $\text{codom}(\varphi_x) = xx^*D$ et $\text{dom}(\varphi_x) = x^*D = x^*xD$ par le lemme 5.18.

- i) Supposons qu'il existe $y, z \in x^*D$ tels que $\varphi_x(y) = \varphi_x(z)$. On a $xy = xz$, soit $xx^*y' = xx^*z'$ en posant $y = x^*y'$ et $z = x^*z'$. On a donc

$$y = x^*y' = x^*xx^*y' = x^*xx^*z' = x^*z' = z$$

et ainsi $\varphi_x \in \mathcal{IP}_D$.

- ii) Avant de montrer que φ est un morphisme, on va montrer que $\varphi_x^{-1} = \varphi_{x^*}$ ³. Par définition, on a

$$\text{dom}(\varphi_x^{-1}) = \text{codom}(\varphi_x) = xx^*D = xD = (x^*)^*D = \text{dom}(\varphi_{x^*})$$

et

$$\text{codom}(\varphi_x^{-1}) = \text{dom}(\varphi_x) = x^*D = x^*xD = \text{codom}(\varphi_{x^*})$$

De plus, pour $a \in \text{dom}(\varphi_x^{-1})$, i.e. $a = xx^*a'$, on a $\varphi_x^{-1}(a) = x^*a' = x^*xx^*a' = x^*a = \varphi_{x^*}(a)$.

- iii) Soient $x, x' \in D$,

$$\begin{aligned} \text{dom}(\varphi_{x'} \circ \varphi_x) &= \varphi_x^{-1}(\text{codom}(\varphi_x) \cap \text{dom}(\varphi_{x'})) \\ &= \varphi_x^{-1}(xx^*D \cap x'^*x'D) \\ &= \varphi_x^{-1}(x'^*x'xx^*D) \\ &= x^*x'^*x'(xx^*D) = x^*x'^*x'(xD) \\ &= (x'x)^*(x'x)D = (x'x)^*D \\ &= \text{dom}(\varphi_{x'x}) \end{aligned}$$

De même, on a

$$\text{codom}(\varphi_x \circ \varphi_{x'}) = \text{codom}(\varphi_{xx'})$$

Maintenant, soit $a \in \text{dom}(\varphi_{xx'})$, $\varphi_{xx'}(a) = xx'a = x(x'a) = (\varphi_x \circ \varphi_{x'})(a)$.

On vient de montrer que φ est un morphisme de semigroupe, et donc de semigroupe inversif.

- iv) Montrons maintenant que φ est injectif. Soient $x, x' \in D$ tels que $\varphi_x = \varphi_{x'}$, cela implique notamment $\text{dom}(\varphi_x) = \text{dom}(\varphi_{x'})$, soit $x^*xD = x'^*x'D$, ce qui entraîne par le lemme 5.18 $x^*x = x'^*x'$. On a donc

$$x = xx^*x = \varphi_x(x^*x) = \varphi_{x'}(x^*x) = \varphi_{x'}(x'^*x') = x'x'^*x' = x'$$



³Cette propriété est pourtant impliquée par la notion de morphisme, mais ici c'est elle qui est primitive.

5.2.4 Algèbre de semigroupes inversif et algèbres stellaires

Dans toute cette section on fixe D un semigroupe inversif.

Définition 5.22 Soit $\mathbb{K}$ un corps, on appelle algèbre du semigroupe D la $\mathbb{K}$ -algèbre $\mathbb{K}[D]$ des combinaisons linéaires finies d'éléments de D à coefficients dans $\mathbb{K}$, i.e. les

$$\sum_{x \in D} \alpha_x x$$

où les $\alpha_x \in \mathbb{K}$ sont presque tous nuls, munie des opérations

$$a(\sum_{x \in D} \alpha_x x) + b(\sum_{y \in D} \beta_y y) = \sum_{x \in D} (a\alpha_x + b\beta_x)x$$

$$(\sum_{x \in D} \alpha_x x)(\sum_{y \in D} \beta_y y) = \sum_{x, y \in D} (\alpha_x \beta_y)xy$$

Rappelons très brièvement la définition d'espace de Hilbert et d'algèbre stellaire.

Définition 5.23 On appelle espace de Hilbert un $\mathbb{C}$ -espace vectoriel muni d'un produit scalaire hermitien et complet vis-à-vis de la norme qu'il induit.

Si H est un espace de Hilbert, on appelle opérateurs les éléments de $\mathcal{L}(H)$ et on les munit de la norme

$$\|f\| = \sup_{x \in H - \{0\}} \frac{\|f(x)\|}{\|x\|}$$

On dit qu'un opérateur est borné si $\|f\| < \infty$, et on note $\mathcal{B}(H)$ l'ensemble des opérateurs bornés.

On appelle adjoint d'un opérateur f l'unique opérateur f^* satisfaisant

$$\forall x, y \quad \langle f(x) | y \rangle = \langle x | f^*(y) \rangle$$

La preuve d'existence et d'unicité est laissé à la discrétion du lecteur, de même que le résultat classique $\|f^*f\| = \|f\|^2$.

La notion d'algèbre stellaire abstrait cette notion d'adjoint.

Définition 5.24 On appelle algèbre stellaire une $\mathbb{C}$ -algèbre munie d'une involution * et d'une norme telle que

- l'algèbre soit de Banach, i.e. elle est normée et complète
- $\forall x$ élément de l'algèbre $\|x^*x\| = \|x\|^2$

Directement, on voit que $\mathcal{B}(H)$ forme une algèbre stellaire, de même pour toute sous-algèbre fermée et close par adjonction. La réciproque est aussi vraie : toute algèbre stellaire se plonge dans un $\mathcal{B}(H)$, ce résultat, appelé construction GNS ⁴, est beaucoup moins directe et nécessiterait (au moins) un chapitre entier à lui tout seul.

Le fait d'avoir noté $\star$ l'opération d'inversion dans les semigroupes n'est pas innocent. En effet, la construction qui va suivre va nous permettre de voir l'algèbre d'un semigroupe inversif comme une algèbre stellaire en superposant l'adjonction et l'inverse.

⁴des noms de ses auteurs : Gelfand, Naimark et Segal.

L'espace $\ell^2(\mathbf{D}) = \{f : \mathbf{D} \rightarrow \mathbb{C}, \sum_{x \in \mathbf{D}} |f(x)|^2 < \infty\}$ est un espace de Hilbert lorsqu'on le muni du produit scalaire

$$\langle a|b \rangle = \sum_{x \in \mathbf{D}} a(x) \overline{b(x)}$$

Il admet comme base les fonctions

$$e_x(y) = \begin{cases} 1 & \text{si } x = y \\ 0 & \text{sinon} \end{cases}$$

Étant donné $f \in \mathcal{IP}_{\mathbf{D}}$ on peut construire un opérateur noté $\mathbf{O}_f$ de $\ell^2(\mathbf{D})$ en posant

$$\mathbf{O}_f(e_x) = e_{f(x)}$$

Soit φ le plongement issu du théorème de Preston-Wagner. A tout $x \in \mathbf{D}$ on associe l'opérateur $\Phi(x)$ sur $\ell^2(\mathbf{D})$ défini par $\Phi_x = \mathbf{O}_{\varphi_x}$. Notons tout d'abord que $\Phi_{xx'} = \Phi_x \circ \Phi'_{x'}$ car $\varphi_{xx'} = \varphi_x \circ \varphi_{x'}$.

Lemme 5.25 Φ_x est un opérateur borné.

Preuve Soit $f \in \ell^2(\mathbf{D})$ telle que $\|f\| = 1$, on a

$$\|\Phi_x(f)\| = \sum_{y \in \mathbf{D}} |f(\varphi_x(y))|^2 \leq \sum_{z \in \text{codom}(\varphi_x)} |f(z)|^2 \leq \sum_{z \in \mathbf{D}} |f(z)|^2 = \|f\| = 1$$

On en déduit $\|\Phi_x\| \leq 1$. ✂

Lemme 5.26 $\Phi_x^* = \Phi_{x^*}$

Preuve On a

$$\begin{aligned} \langle \Phi_x(f)|g \rangle &= \sum_{y \in \mathbf{D}} \Phi_x(f)(y) \overline{g(y)} \\ &= \sum_{y \in \text{dom}(\varphi_x)} f(\varphi_y) \overline{g(y)} \\ &= \sum_{z \in \text{codom}(\varphi_x)} f(z) \overline{g(\varphi_{x^*}(z))} \\ &= \overline{\langle f|\Phi_{x^*}(g) \rangle} \end{aligned}$$
✂

On étend Φ à $\mathbb{C}[\mathbf{D}]$ en posant $\Phi_{\sum_x \alpha_x x} = \sum_x \alpha_x \Phi_x$.

On a donc $\Phi_{\sum_x \alpha_x x}(e_y) = \sum_x \alpha_x e_{\varphi_x(y)}$

Φ réalise un morphisme d'algèbre stellaire, i.e. morphisme d'algèbre et préservation de $\star$. L'injectivité de φ assure que $\Phi_x \neq \Phi_y$ pour $x \neq y \in \mathbf{D}$, cependant cette injectivité ne s'étend pas nécessairement à $\mathbb{C}[\mathbf{D}]$.

Lemme 5.27 Si $\mathbf{D}$ est régulier à gauche alors Φ appliqué à $\mathbb{C}[\mathbf{D}]$ est injectif

Preuve Soit $\bar{x} = \sum_x \alpha_x x \in \mathbb{C}[\mathbf{D}]$ tel que $\Phi_{\bar{x}} = 0$. On a donc, $\forall y, \sum_x \alpha_x e_{\varphi_x(y)} = \sum_x \alpha_x e_{xy} = 0$, les e_z étant libres, on en déduit : $\forall y, z \in \mathbf{D}, \sum_{x \in \mathbf{D}, xy=z} \alpha_x = 0$.

Or l'hypothèse nous assure que cette somme ne comporte qu'un seul terme, et donc $\alpha_x = 0$. ✂

On a donc le théorème suivant :

Théorème 5.28 $\mathbb{C}[\mathbf{D}]$ est une algèbre pré-stellaire.

Par *pré-stellaire* on entend que tous les axiomes des algèbres stellaires sont validés excepté la complétude vis-à-vis de la norme.

Preuve On vient de montrer que Φ était un morphisme d'algèbre involutive de $\mathbb{C}[\mathbf{D}]$ dans $\mathcal{B}(\ell^2(\mathbf{D}))$. On peut faire remonter la norme pour munir $\mathbb{C}[\mathbf{D}]$ d'une norme d'algèbre. Notons que l'image de $\mathbb{C}[\mathbf{D}]$ n'est pas nécessairement fermée, et donc l'algèbre n'est pas forcément de Banach. $\blacktriangleleft$

Remarque 5.29 Notre construction n'est pas la plus naturelle lorsque notre semigroupe de départ est en fait un monoïde inversif possédant un zéro. Si on procède directement, on est amené à considérer un élément e_0 dans $\ell^2(\mathbf{D})$, et de fait, $\Phi_0 = \text{proj}_{e_0} \neq 0$. Cependant, $\ell^2(\mathbf{D}) = \text{Vect}(e_0) \perp H$ et $\Phi_0|_H = 0$. On note $\mathbb{C}[\mathbf{D}]_0$ l'algèbre de $\mathbf{D}$ associée à H , c'est-à-dire celle où l'on a superposé le 0 de $\mathbf{D}$ à celui de $\mathbb{C}$.

5.2.5 Construction explicite de l'algèbre de MLL

Dans cette section on fixe un monoïde inversif avec zéro généré par deux éléments p et q , satisfaisant

$$p^\star p = q^\star q = 1 \quad p^\star q = q^\star p = 0$$

En orientant les relations de la gauche vers la droite, on constate tout de suite que chaque élément de $\mathbf{D} - \{0\}$ peut-être mis sous la forme $ab^\star$, où $a, b \in \{p, q\}^\star$.

φ_p est de domaine plein, car $\text{dom}(\varphi) = p^\star \mathbf{D} = p^\star p \mathbf{D} = \mathbf{D}$, et on a $\varphi_p(ab^\star) = pab^\star$. $\text{dom}(\varphi_{p^\star}) = p\mathbf{D}$ et $\varphi_{p^\star}(pab^\star) = ab^\star$.

En utilisant un codage binaire direct, on peut définir une bijection n de $\{p, q\}^\star$ dans $\mathbb{N}$:

$$n(pw) = 2n(w) \quad n(qw) = 2n(w) + 1 \quad n(1) = 0$$

n peut être prolongée de $\mathbf{D} - \{0\}$ vers $\mathbb{Z}$ en posant $n(ab^\star) = n(a) - n(b)$, mais bien sur il cesse d'être injectif.

On remarque tout de suite que $\ell^2(\{p, q\}^\star) \simeq \ell^2(\mathbb{N})$, en posant $e_w \simeq e_{n(w)}$, et en restreignant notre construction précédente à ce sous-espace, on obtient :

$$\Phi_p(e_k) = e_{2k} \quad \Phi_q(e_k) = e_{2k+1}$$

Cette construction redonne donc exactement la définition en terme d'opérateurs telle qu'elle est introduite dans [Gir89a].

Chapitre 6

Chemins et réduction

6.1 Graphes de réseaux

- Soit $r = (\sigma_w, \sigma_c)$ un réseau d'interaction, on construit deux graphes sur les ports de r :
- le graphe non orienté des fils $G_w(r) = (P(r), E_w)$ défini par la relation binaire symétrique : $pE_w p' \iff p = \sigma_w(p')$
 - le graphe orienté des cellules $G_c(r) = (P(r), E_c)$ où

$$E_c = \{(\text{pax}_i(c), \text{pal}(c)), (\text{pal}(c), \text{pax}_i(c)) \mid c \text{ cellule de } r, 1 \leq i < |c|\}$$

L'orientation du graphe $G_c(r)$ peut sembler superflue, cependant des deux arêtes $(\text{pax}_i(c), \text{pal}(c))$ et $(\text{pal}(c), \text{pax}_i(c))$ nous considérerons la première comme l'orientation naturelle. Afin d'accen-
tuer ce choix, nous notons c_i cette première arête et nous considérons $G_c(r)$ comme engendré par les c_i et clos par inversion de l'orientation. On note c_i^r l'inversion de c_i .

La figure 6.1 donne une représentation visuelle de ces graphes.

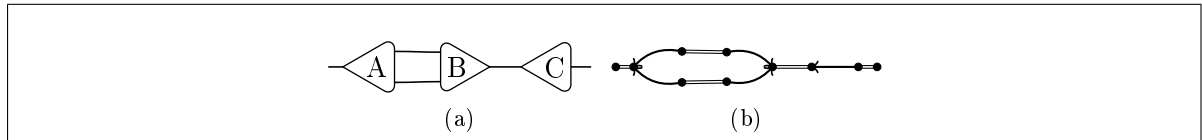


Figure 6.1: a) un réseau d'interaction r et b) ses graphes associés : $G_w(r)$ est représenté en traits doubles et les arêtes principales, les c_i , de $G_c(r)$ en traits pleins

Fait 6.1 Soit r un réseau d'interaction. Les graphes $G_w(r)$ et $G_c(r)$ sont simples, c'est-à-dire qu'ils ne contiennent ni boucles ni arêtes parallèles.

6.2 Chemins

6.2.1 Définition

Définition 6.2 Soit r un réseau d'interaction, un chemin dans r est une suite finie

$$p_1, e_1, p_2, \dots, e_n, p_{n+1}$$

telle que

- $p_i \in P(r)$
- soit les e_{2i} sont des arêtes de $G_w(r)$ et les e_{2i+1} arêtes de $G_c(r)$, soit l'inverse
- e_i relie le port p_i au port p_{i+1}

p_1 est appelé la source du chemin et p_{n+1} son but. On note id_p le chemin vide de source et but p .

6.2.2 Catégorie des chemins

Les chemins dans un réseau ne sont pas directement des chemins dans un graphe : ils alternent entre le graphe $G_w(r)$ et le graphe $G_c(r)$. De ce fait, deux chemins dont la source de l'un est le but de l'autre, ne sont pas nécessairement composables. On ne peut donc pas parler de petite catégorie des chemins dans un réseau d'interaction. Cependant on peut y remédier en considérant des sommes formelles de chemins, on note 0 la somme vide et on pose, pour φ chemin de p' à p'' et ψ chemin de p à p' ,

$$\varphi\psi = \begin{cases} \varphi\psi & \text{si c'est un chemin valide} \\ 0 & \text{sinon} \end{cases}$$

qui est un chemin de p à p'' . On peut maintenant définir la catégorie $\mathbf{Path}_r$ dont les objets sont les ports de r et les morphismes les somme formelles de chemins muni de la nouvelle composition.

6.2.3 Représentation usuelle

La restriction sur la composition des arêtes permet d'omettre les ports internes dans la description d'un chemin, de même si le chemin compte au moins une arête de $G_c(r)$ alors sa source et son but sont uniquement déterminé. Ainsi, on notera juste $e_1 \dots e_n$ le chemin $p_1, e_1, \dots, p_{n+1}$. On peut encore plus simplifier la représentation d'un chemin en remarquant qu'étant donné deux arêtes de $G_c(r)$ il n'existe au plus qu'une arête de $G_w(r)$ liant le but de l'une à la source de l'autre. On peut donc se passer d'indiquer les fils compris entre deux arêtes de $G_c(r)$ dans la description d'un chemin.

On note $\mathfrak{P}(r)$ l'ensemble des chemins de r , et $\forall p, p' \in P(r)$, on note $\mathfrak{P}_{p \rightarrow p'}(r)$ l'ensemble des chemins de source p et de but p' et on note $\mathfrak{P}_f(r)$ l'ensemble des chemins reliant deux ports libres de r . On note $\mathfrak{P}(r)^+$ l'ensemble des sommes formelles de chemins dans r , les autres ensembles étant notés de manière similaire. On a notamment

$$\mathfrak{P}_{p \rightarrow p'}(r)^+ = \mathbf{Path}_r(p, p')$$

Remarquons que $\mathfrak{P}(r)$ peut être infini comme en témoigne le réseau suivant :



On étend l'opération d'inversion d'orientation en posant

$$(e_1 \dots e_n)^r = e_n^r \dots e_1^r \text{ et } w^r = w \ \forall w \in E_w(r)$$

Cette opération réalise une bijection de $\mathfrak{P}_{p \rightarrow p'}(r)$ dans $\mathfrak{P}_{p' \rightarrow p}(r)$.

6.2.4 Chemins dans des boîtes

La définition précédente s'applique directement et nous permet de définir les chemins dans un réseau à boîtes, en effet, une boîte est avant tout une cellule. Cependant, on souhaiterait pouvoir effectuer un chemin *en profondeur* en allant dans les boîtes continuer le chemin.

Prenons un réseau R contenant une boîte nommée c , elle même contenant un réseau sans boîtes r . Les chemins dans r sont donnés par la définition précédente. Soient p, p' deux ports de c , dont on note p_r, p'_r les ports associés dans $P_f(r)$, on connaît les chemins de p_r à p'_r , à chacun de ces chemins on peut en quelque sorte associer une *traversée* de c allant de p à p' .

Ainsi, à la place de considérer, dans ce cas, $G_c(R)$ comme engendré par les arêtes c_i , on va rajouter des arêtes noté $c(\varphi)$ pour φ chemin de r , reliant naturellement les ports associés. Bien que $G_c(R)$ reste symétrique, il n'est pas possible de définir une orientation naturelle permettant de l'engendrer par symétrisation comme précédemment. En effet, un chemin allant d'un port à lui-même n'a pas une orientation privilégiée.

Cette redéfinition de $G_c(R)$ peut être prolongée aux réseaux à boîtes sans restriction sur leur contenu. Notons tout de même que le fait 6.1 n'est plus valide dans la mesure où deux arêtes de la forme $c(\varphi)$ peuvent avoir même source et même but.

6.3 Réduction de chemins

6.3.1 Pseudo-chemins

Pour définir la réduction des chemins, on va devoir passer par une notion intermédiaire.

Définition 6.3 Soient R_1, R_2 deux réseaux d'interactions à supports disjoints et f une injection partielle de $P_f(R_1)$ dans $P_f(R_2)$. On appelle pseudo-chemin de $R_1 \overset{f}{\rightsquigarrow} R_2$ une suite $(\varphi_1, \dots, \varphi_n)$ de chemins alternant entre R_1 et R_2 et telle que, si l'on note p_i (resp. p'_i) la source (resp. le but) de φ_i ,

- $p_1, p'_n \notin \text{dom}(f) \cup \text{codom}(f)$
- pour tout $i < n$, si φ_i est un chemin de R_1 (resp. de R_2) alors $p_{i+1} = f(p'_i)$ (resp. $p_{i+1} = f^{-1}(p'_i)$).

Un pseudo-chemin est donc un chemin de $R_1 \overset{f}{\rightsquigarrow} R_2$ allant du port p_1 au port p'_n , au recollement près de ses fils successifs. En effet, soit $w_1 \dots w_m$ une succession de fils présente dans un pseudo-chemin φ , et maximale pour l'inclusion. La première condition sur le pseudo-chemin nous assure que la source de w_1 noté p , ainsi que le but de w_m noté p' , appartiennent au $\text{dom}(\sigma_w \overset{f}{\rightsquigarrow} \tau_w)$, pour $R_1 = (\sigma_w, \sigma_c)$ et $R_2 = (\tau_w, \tau_c)$. De plus, la seconde condition nous assure que $(\sigma_w \overset{f}{\rightsquigarrow} \tau_w)(p) = p'$.

Définition 6.4 Dans les conditions de la définition 6.3, soit φ un pseudo-chemin, on note $\chi(\varphi)$ la suite de chemins obtenue en substituant à toute succession de fils dans φ , maximale pour l'inclusion, le fil (p, p') de $R_1 \overset{f}{\rightsquigarrow} R_2$, où p (resp. p') est la source (resp. le but) de la succession de fils.

Cette définition est correcte du fait de la remarque précédente, et l'on a de plus :

Fait 6.5 Dans les conditions de la définition précédente,

$$\chi(\varphi) \in \mathfrak{P}(R_1 \xleftrightarrow{f} R_2)$$

Preuve Seule la condition d'alternance est à vérifier. Or, un pseudo-chemin alterne entre successions de fils maximales et arêtes de cellules, et χ réduisant les successions de fils maximales vers de simples fils, on est assuré que $\chi(\varphi)$ soit un chemin. $\blacktriangleright$

6.3.2 Construction de la réduction

Soit $\mathcal{R} = (R_r^{I_r}, R_p^{I_p})$ une règle d'interaction pour (s_1, s_2) . Et soit $R_0 = R^I \xleftrightarrow{\alpha} \alpha(R_r)^{\alpha(I_r)}$ un réseau d'interaction sur lequel s'applique la règle. On note c_1 (resp. c_2) l'image par α de la cellule de symbole s_1 (resp. s_2) dans R_r . On fixe une réduction $R_0 \rightarrow_{\mathcal{R}} R^I \xleftrightarrow{\beta} \beta(R_p)^{\alpha(I_p)} = R_1$.

On dit qu'un chemin est *assez long* vis-à-vis de cette réduction, lorsque ni sa source ni son but n'appartient à c_1 ou c_2 . On note $\mathfrak{P}_{\mathcal{R}, c_1, c_2}(R_0)$, ou juste $\mathfrak{P}_{\mathcal{R}}(R_0)$, l'ensemble de ces chemins. Un chemin assez long est donc un chemin qui traverse toujours complètement un redex.

Les traversées de redex par un chemin assez long ont la forme

$$wc_{a, i_1} c_{b, j_1}^r \dots c_{a, i_{2n}} c_{b, j_{2n}}^r w'$$

ou

$$wc_{a, i_1} c_{b, j_1}^r \dots c_{b, i_{2n+1}} c_{a, j_{2n+1}}^r w'$$

pour $a \neq b \in \{1, 2\}$ et $1 \leq i_k < |c_a|$, $1 \leq j_k < |c_b|$. Un sous-chemin de la forme $c_{a, i} c_{b, j}^r$, avec $a \neq b \in \{1, 2\}$, est appelé *traversée de redex simple*.

La figure 6.2 présente différentes traversées de redex.

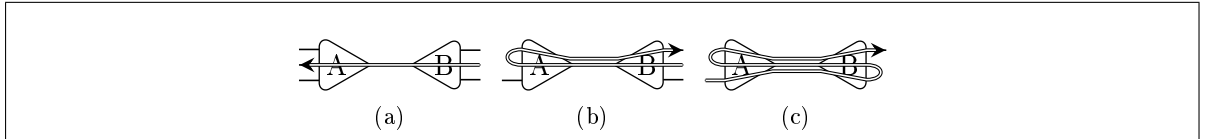


Figure 6.2: Traversées de redex : a) simple, b) avec un nombre de rebond pair, c) impair

On va construire une application $\delta_{\mathcal{R}} : \mathfrak{P}_{\mathcal{R}}(R_0)^+ \rightarrow \mathfrak{P}(R_1)^+$ associant à un chemin la somme de ses réduits.

Il est immédiat qu'un chemin ne traversant pas le redex est préservé par réduction et il suffit donc de définir l'image des traversées de redex. De plus, l'image de ces traversées dépend des chemins du motif et des éventuels fils *rebonds* liant deux portes auxiliaires d'une des cellules. Pour les premiers, remarquons que l'image par réduction de $c_{1, i} c_{2, j}^r$ est, si l'on pose $p_{1, i} = f^{-1}(1, i)$ et $p_{2, j} = f^{-1}(2, j)$,

$$\mathfrak{P}_{p_{1, i} \rightarrow p_{2, j}}(r) = \begin{cases} \emptyset \\ \{(p_{1, i}, p_{2, j})\} \in G_w(r) \\ \{\varphi_1, \dots, \varphi_n\} \quad \text{où } \forall i, \varphi_i \cap G_c(r) \neq \emptyset \end{cases} \quad (6.1)$$

Le premier cas correspond à $p_{1, i}$ et $p_{2, j}$ déconnectés, le second à avoir ces deux ports connectés par un fil et dans le cas restant plusieurs possibilités de connections, toutes passant à travers au moins une cellule.

On pose $\Phi(1, i \mapsto 2, j) = \sum_{\phi \in \mathfrak{P}_{p_{1,i} \mapsto p_{2,j}}(r)} \phi$. Soit $\varphi = w(\prod_k c_{1,i_k} c_{2,j_k}^x)w$ une traversée de redex, les autres cas se ramenant trivialement à celui-là, on aimerait pouvoir poser $\delta_{\mathcal{R}}(\varphi) = w(\prod_k \Phi(1, i_k \mapsto 2, j_k))w'$, or c'est une somme de pseudo-chemins, on se ramène à des chemins à l'aide de χ , qu'on a trivialement étendue aux sommes. Ainsi la définition réelle de δ est

$$\delta_{\mathcal{R}}(\varphi) = \chi(w(\prod_k \Phi(1, i_k \mapsto 2, j_k))w')$$

Le même raisonnement sur la composition par exécution des fils permet de déduire de la propriété 2.7 le théorème suivant :

Théorème 6.6 (Confluence forte de la réduction de chemins) *Soient $\mathcal{R}_1$ et $\mathcal{R}_2$ deux règles d'interactions et $\delta_{\mathcal{R}_1}, \delta_{\mathcal{R}_2}$ les deux réductions de chemins correspondant à l'application de ces règles sur un réseau d'interaction R . On a*

$$\delta_{\mathcal{R}_1} \circ \delta_{\mathcal{R}_2} = \delta_{\mathcal{R}_2} \circ \delta_{\mathcal{R}_1}$$

Remarque 6.7 *Dans l'énoncé du théorème précédent, il est implicite, pour que l'égalité fasse sens, qu'elle s'applique sur des chemins assez long vis-à-vis des deux instances considérées. Un moyen simple de s'affranchir de ces conditions et de ne considérer que des chemins de port libre à port libre. En effet, de tels chemins, maximaux pour l'inclusion, sont garantis d'être assez long vis-à-vis de n'importe quelle réduction.*

Notons le fait suivant permettant de relever des chemins port libre à port libre.

Fait 6.8 *Soit $R \rightarrow_{\mathcal{R}} R'$ une réduction. On a*

$$\forall \varphi' \in \mathfrak{P}_f(R'), \exists ! \varphi \in \mathfrak{P}_f(R), \delta_{\mathcal{R}}(\varphi) = \varphi' + \sum_i \varphi_i$$

Preuve C'est une implication directe du fait qu'à chaque traversée par φ' du réseau qui a remplacé le redex correspond une unique traversée simple du redex. $\blacktriangleright$

6.3.3 Possibilité d'une définition alternative plus simple

Au vu de la définition des chemins et de leur réduction, ce qui a été fait ici peut sembler anormalement compliqué. On va donner ici différents contre-exemples à des définitions plus simples.

6.3.3.a Chemins sans fils

Une des premières remarques faites après avoir défini les chemins dans toute leur généralité a été de dire que les fils sont le plus souvent superflus. En effet entre deux arêtes de cellules, il n'existe au plus qu'un fil les liant. On pourrait alors être tenté de définir une structure de monoïde avec zéro engendré par E_c et tel que $ee' = 0$ pour $e, e' \in E_c$ si il n'existe pas de fil liant la source de e' au but de e . Les avantages d'une telle structure algébrique sont évidents. Malheureusement, il n'est pas possible de définir un chemin ne passant que par un fil, et de tels chemins sont cruciaux lorsqu'ils lient deux ports libres.

6.3.3.b Traversées de redex simples

Les traversées de redex comme elles ont été exposées sont complexes, et de ce fait, oblige à effectuer une difficile définition de $\delta_{\mathcal{R}}$. On pourrait pour pallier cela définir un chemin comme assez long dès qu'il n'a pas comme source ou but une des extrémité du fil actif : en effet, les ports auxiliaires étant *recollés* sur les ports libres du motif, il est possible de s'y retrouver. Là encore, il s'agit d'une fausse bonne idée, et il suffit de considérer le réseau suivant



se réduisant en



où le chemin dessiné ne peut pas être réduit : il devient un *bout de fil*. Intuitivement, on comprend alors pourquoi il n'est pas *assez long*.

6.3.4 Réduction de chemins et boîtes

La réduction de chemins dans les réseaux d'interaction avec boîtes est sensiblement plus complexe que le cas précédent. En effet, un chemin dans une boîte ne passe pas nécessairement par sa porte principale, il devient alors difficile d'exprimer une règle locale de réduction autour des traversées du fil actif étant réduit.

De ce fait, nous n'allons pas donner de définitions générales pour toute règle faisant intervenir des boîtes. A la place, nous donnons la définition de la réduction de chemins dans le cadre des opérations définies au paragraphe 3.2.

Pour décrire complètement la réduction de chemins, il nous suffit de décrire la réduction d'une traversée du redex. En vertu des définitions précédentes, on peut se ramener à l'étude de traversées simples, pourvu que l'on réduise les chemins à l'aide de χ . Les traversées de redex sont de trois natures différentes ici, selon comment s'effectue la traversée de la boîte réduite :

- les chemins *pal/pal* effectuant un rebond de la porte principale vers elle-même à l'intérieur de la boîte ;
- les chemins *pal/pax* traversant la boîte depuis la porte principale vers une porte auxiliaire, cela correspond au cas usuel des réseaux sans boîtes ;
- les chemins *pax/pax* traversant la boîte d'une porte auxiliaire vers une autre.

6.3.4.a Déballage

On considère un redex pour la règle de déballage, soient c la cellule de symbole d_s et c' la s -boîte. Si φ est un chemin

pal/pax on pose $\delta(c_1 c'(\varphi)) = \varphi$

pal/pal on pose $\delta(c_1 c'(\varphi) c_1^r) = \varphi$

pax/pax on pose $\delta(c'(\varphi)) = \varphi$

6.3.4.b Duplication

On considère maintenant un redex pour la règle de duplication, la cellule de duplication est notée c , la boîte initiale b et ses deux réduits b^1 et b^2 . On note c^i la cellule de duplication reliant les i -èmes ports auxiliaires de b^1 et b^2 . Si φ est un chemin

pal/pax on pose $\delta(c_i b(\varphi)) = b^i(\varphi) c_i^k$, $\forall i \in \{1, 2\}$ et où k est la porte auxiliaire but de φ

pal/pal on pose $\delta(c_i b(\varphi) c_j) = \begin{cases} 0 & \text{si } i \neq j \\ b^i(\varphi) & \text{sinon} \end{cases}$

pax/pax on pose, si φ va de la porte auxiliaire i à la porte auxiliaire j ,

$$\delta(b(\varphi)) = \sum_{k \in \{1, 2\}} c_k^{i_r} b^k(\varphi) c_k^j$$

6.3.4.c Effacement

Les seuls chemins pouvant être réduit dans un redex de la règle d'effacement sont les chemins **pax/pax**, et dans ce cas on a $\delta(b(\varphi)) = 0$.

6.3.4.d Enfouissement

On considère une boîte b' s'enfouissant dans une boîte b , devenant ainsi une boîte b_r . On a $\delta(b(\varphi) b'(\varphi')) = b_r(\varphi b'(\varphi'))$ et si φ est un chemin **pax/pax** de b' on a $\delta(b'(\varphi)) = b_r(b'(\varphi))$.

Remarque 6.9 *Les opérations de déballage, duplication et effacement préservent la nature des chemins dans une boîte. C'est-à-dire que le réduit d'un chemin **pax/pax** est toujours un chemin **pax/pax**, de même pour les autres types de chemins. Ce n'est pas le cas de l'opération d'enfouissement : un chemin de la forme $b(\varphi) b'(\varphi') b(\varphi'')$, où, nécessairement pour que la composition fasse sens, φ est **pal/pax**, φ' est **pal/pal** et φ'' est **pax/pal**, sera réduit sur $b_r(\varphi b'(\varphi') \varphi'')$ où le chemin dans b_r est un chemin **pal/pal**.*

6.3.4.e Confluence

On peut démontrer, voir [Reg92], le résultat suivant :

Théorème 6.10 *La réduction de chemins dans les réseaux d'interactions avec boîtes, munis du déballage, de la duplication, de l'effacement et de l'enfouissement, est confluente.*

Corollaire 6.11 *La réduction de chemins de MELL est confluente.*

Chapitre 7

Géométrie de l'Interaction

7.1 Persistance et pondérations fidèles

7.1.1 Dans les réseaux d'interaction

7.1.1.a Persistance

Le théorème 6.6 nous permet de donner sans ambiguïté la définition suivante :

Définition 7.1 Soient $\mathcal{B}$ une bibliothèque et R un réseau d'interaction tel qu'il existe une chaîne de réductions

$$R \rightarrow_{\mathcal{R}_1} R_1 \cdots \rightarrow_{\mathcal{R}_n} R_n$$

où les $\mathcal{R}_i \in \mathcal{B}$ et R_n est en forme normale, i.e. R est normalisable, on dit qu'un chemin $\varphi \in \mathfrak{P}_f(R)$ est persistant pour $\mathcal{B} \iff$

$$\delta_{\mathcal{R}_n} \circ \cdots \circ \delta_{\mathcal{R}_1}(\varphi) \neq 0$$

Notons que le choix de ne parler que de chemins dans $\mathfrak{P}_f(R)$ est issu de la remarque 6.7. En effet, de tels chemins sont la garantie que la réduction le long d'une chaîne quelconque soit bien définie.

Remarque 7.2 La persistance n'a pas besoin de la confluence forte pour être définie, seule l'égalité de la réduction de chemins le long de deux chaînes de réduction vers la forme normale suffit. Or, cette égalité est entraînée par la confluence seule.

7.1.1.b Pondération

Définition 7.3 Soit R un réseau d'interaction, on appelle c -pondération, ou simplement pondération, une fonction

$$w : E_c(R) \rightarrow D$$

où D est un semigroupe inversif, et telle que $w(e^r) = w(e)^\star$ pour tout $e \in E_c(R)$.

Définition 7.4 Une pondération w est dite générique si $w(c_i)$ ne dépend que de i et du symbole de c .

Définition 7.5 Soit $\mathcal{S}' \subseteq \mathcal{S}$, il suffit de donner $w(s, i)$ pour $s \in \mathcal{S}'$ et $1 \leq i < \alpha(s)$ pour pouvoir définir w une pondération générique de tout réseau de $\mathfrak{R}(\mathcal{S}')$.

Dans ce cas, w est dite $\mathcal{S}'$ -générique.

Toute pondération w peut être étendue à $\mathfrak{P}(R)$ en posant $w(w) = 1$ pour $w \in E_w(R)$, $w(id_p) = 1$ et $w(e\varphi) = w(\varphi)w(e)$. De plus si D possède un zéro, on peut étendre w de $\mathfrak{P}(R)^+$ dans $\mathbb{C}[D]_0$ ¹ en posant $w(0) = 0$ et $w(\sum \varphi_i) = \sum w(\varphi_i)$.

Définition 7.6 Soit R un réseau d'interaction et w une pondération vers un semigroupe avec zéro, on dit qu'un chemin φ de R est régulier vis-à-vis de w si $w(\varphi) \neq 0$.

w est dite fidèle pour $\mathcal{B}$, une bibliothèque, si (φ persistant pour $\mathcal{B} \iff \varphi$ régulier vis-à-vis de w).

On pose $\mathfrak{Pr}_f(R, w) = \{\phi \in \mathfrak{P}_f(R) \mid \phi \text{ régulier vis-à-vis de } w\}$, on notera juste $\mathfrak{Pr}_f(R)$ lorsqu'il n'y aura pas d'ambiguïté sur la pondération considérée.

Définition 7.7 Soit $\mathcal{B}$ une bibliothèque close pour $\mathcal{S}'$, $\mathbf{R} \subseteq \mathfrak{R}(\mathcal{S}')$ et w une pondération $\mathcal{S}'$ -générique à valeur dans D .

Si $D = \langle \bar{g} \mid \bar{r} \rangle$ et $\forall s \in \mathcal{S}', 1 \leq i < \alpha(s), \exists k, w(s, i) = g_k$ ($\neq 0$ si D en a un) alors w est dite simple².

w est dite standard pour $\mathcal{B}$ et $\mathbf{R}$ si pour tout réseau $R \in \mathbf{R}$, normal pour $\rightarrow_{\mathcal{B}}$, et tout chemin $\varphi \in \mathfrak{P}(R)$ on a $w(\varphi) \neq 0$

w est dite fortement fidèle pour $\mathcal{B}$ et $\mathbf{R}$ si elle est standard et que, de plus, pour tout $R \in \mathbf{R}$ tel qu'il existe $\mathcal{R} \in \mathcal{B}$ avec $R \rightarrow_{\mathcal{R}} R'$, on a

$$\forall \varphi \in \mathfrak{P}_f(R) \quad w(\varphi) = w(\delta_{\mathcal{R}}(\varphi))$$

Propriété 7.8 Soit $\mathcal{B}$ une bibliothèque close pour $\mathcal{S}'$ $\mathbf{R} \subseteq \mathfrak{R}(\mathcal{S}')$ et w une pondération $\mathcal{S}'$ -générique.

1. w fortement fidèle pour $\mathcal{B}$ entraîne w fidèle pour $\mathcal{B}$.
2. On suppose donné un typage associé à $\mathcal{S}'$ tel que $\mathcal{B}$ soit complète vis-à-vis de celui-ci, on note $\mathfrak{R}(\mathcal{S}')^t$ l'ensemble des réseaux bien typés. Alors D normal et w simple entraîne w standard pour $\mathcal{B}$ et $\mathfrak{R}(\mathcal{S}')^t$.
3. Si de plus $\rightarrow_{\mathcal{B}}$ normalise sur $\mathfrak{R}(\mathcal{S}')^t$ et w est fortement fidèle, alors les poids des chemins réguliers de $\mathfrak{R}(\mathcal{S}')^t$ peuvent être mis sous la forme $\sum_i a_i b_i^*$.

Preuve 1) Soit $R \in \mathbf{R}$ tel que

$$R \rightarrow_{\mathcal{R}_1} R_1 \cdots \rightarrow_{\mathcal{R}_n} R_n$$

où R_n est en forme normale.

Soit $\varphi \in \mathfrak{P}(R)$, on pose

$$\delta_{\mathcal{R}_n} \circ \cdots \circ \delta_{\mathcal{R}_1}(\varphi) = \psi$$

En vertu de la fidélité forte de w on a $w(\varphi) = w(\psi)$.

Or w standard entraîne $w(\psi) = 0 \iff \psi = 0$ et on a donc φ persistant ssi il est régulier.

¹cf. la remarque 5.29

²Cette définition dépend de la présentation, en pratique nous ne regarderons jamais simultanément deux présentations d'un même semigroupe.

2) C'est une conséquence de la propriété 4.4. En effet, un chemin dans un réseau en forme normale absolue est de la forme $\prod_i c'_{i,l_i} \prod_j c_{j,k_j}$ et son image par w simple est de la forme $\prod_j g_j (\prod_i g_i)^*$ qui est $\neq 0$ car D est normal.

3) Si $\rightarrow_{\mathcal{B}}$ normalise, soit R réseau bien typé, alors il existe R_0 forme normale de R . L'argument précédent s'applique dans R_0 et montre que tout chemin y a un poids de la forme ab^* . Or, si $\varphi \in \mathfrak{P}(R)$, on a $w(\varphi) = w(\delta(\varphi)) = \sum_i w(\varphi_i)$ avec $\varphi_i \in \mathfrak{P}(R_0)$, et donc ayant un poids de la forme ab^* . $\blacktriangleright$

Dans la mesure où la pondération d'un fil vaut toujours 1, pour montrer qu'une pondération est fortement fidèle il suffit de le montrer pour toute règle $\mathcal{R}$ pour (s_1, s_2) dans $\mathcal{B}$, sur des chemins de la forme $c_{1,i}c_{2,j}^r$ avec c_i de symbole s_i . Si l'on garde les notations de l'équation 6.1, il suffit donc d'avoir

$$w(c_{1,i}c_{2,j}^r) = \begin{cases} 0 & \text{si } p_{1,i} \text{ et } p_{2,j} \text{ sont deconnectés} \\ 1 & \text{si } p_{1,i} \text{ et } p_{2,j} \text{ sont connectés par un fil} \\ \sum_i w(\varphi_i) & \text{sinon} \end{cases} \quad (7.1)$$

pour affirmer que w est fortement fidèle.

7.1.2 Dans les réseaux d'interaction avec boîtes

En vertu de la remarque 7.2 et du théorème 6.10, on conserve la définition de la persistance. Pour définir la notion de pondération appropriée, il est plus simple de définir directement la notion de pondération générique. En effet, il est naturel que le poids d'un chemin passant à travers une boîte dépende du poids du chemin dans la boîte.

Définition 7.9 Soit $\mathfrak{R}(\mathcal{S}', \mathcal{S}_b)$ un ensemble de réseaux d'interaction dont les cellules sont étiquetées par $\mathcal{S}'$ et les boîtes par $\mathcal{S}_b$. On appelle pondération $(\mathcal{S}', \mathcal{S}_b)$ -générique la donnée d'un semigroupe D ,

- d'éléments $w(s, i)$ pour $s \in \mathcal{S}'$ et $1 \leq i < \alpha(s)$,
- et d'applications $!_{ll}^s, !_{lx}^s$ et $!_{xx}^s$ de D dans D pour tout symbole s de boîtes.

Pour toute w pondération $(\mathcal{S}', \mathcal{S}_b)$ -générique et $R \in \mathfrak{R}(\mathcal{S}', \mathcal{S}_b)$, on définit une pondération w sur R par

- $w(c_i) = w(s, i)$ si s est le symbole de c
-

$$w(c(\varphi)) = \begin{cases} !_{ll}^s(w(\varphi)) & \text{si } \varphi \text{ est un chemin pal/pal} \\ !_{lx}^s(w(\varphi)) & \text{si } \varphi \text{ est un chemin pal/pax} \\ !_{lx}^s(w(\varphi)^*)^* & \text{si } \varphi \text{ est un chemin pax/pal} \\ !_{xx}^s(w(\varphi)) & \text{si } \varphi \text{ est un chemin pax/pax} \end{cases}$$

où c est une s -boîte.

On peut définir de même les notions de fidélité et de fidélité forte. Cependant il ne suffit plus de vérifier les équations 7.1 pour prouver la fidélité forte. En effet, il faut tenir compte au cas par cas des différentes opérations pouvant être effectuées sur les boîtes.

On va considérer des s -boîtes et on notera juste $!_l(x) = !_{ll}^s(x)$, de même pour $!_{lx}(x)$ et $!_{xx}(x)$.

déballage Soit c une s -boîte réduite par déballage par une cellule c' de symbole d , et soit φ un chemin

pal/pal

$$\mathbf{w}(d, 1)^*!_l(\mathbf{w}(\varphi))\mathbf{w}(d, 1) = \mathbf{w}(\varphi) \quad (7.2)$$

pal/pax

$$!_{lx}(\mathbf{w}(\varphi))\mathbf{w}(d, 1) = \mathbf{w}(\varphi) \quad (7.3)$$

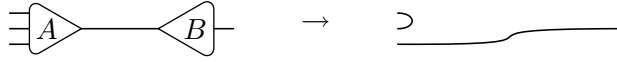
pax/pax

$$!_{xx}(\mathbf{w}(\varphi)) = \mathbf{w}(\varphi) \quad (7.4)$$

L'équation 7.4 est un sérieux frein à l'obtention d'une pondération fortement persistante. En effet, elle force $!_{xx}(x) = x$ pour tout x . Ce qui a pour effet d'annuler des chemins persistants, comme le chemin représenté dans le réseau suivant :



muni d'une réduction détruisant ce chemin en l'absence de la boîte, comme par exemple



duplication Considérons une s -boîte c dupliquée par une cellule de duplication c' de symbole δ , et soit φ un chemin

pax/pax on doit avoir

$$!_{xx}(\mathbf{w}(\varphi)) = \begin{aligned} & \mathbf{w}(\delta, 1)!_{xx}(\mathbf{w}(\varphi))\mathbf{w}(\delta, 1)^* \\ & + \mathbf{w}(\delta, 2)!_{xx}(\mathbf{w}(\varphi))\mathbf{w}(\delta, 2)^* \end{aligned} \quad (7.5)$$

pal/pal

$$\mathbf{w}(\delta, i)^*!_l(\mathbf{w}(\varphi))\mathbf{w}(\delta, j) = \begin{cases} 0 & \text{si } i \neq j \\ !_l(\mathbf{w}(\varphi)) & \text{sinon} \end{cases} \quad (7.6)$$

pal/pax

$$!_{lx}(\mathbf{w}(\varphi))\mathbf{w}(\delta, i) = \mathbf{w}(\delta, i)!_{lx}(\mathbf{w}(\varphi)) \quad (7.7)$$

L'équation 7.7 exprime la commutation entre l'image de $!_{lx}()$ et les $\mathbf{w}(\delta, i)$. Si l'on a également de telles commutations avec $!_l()$ et $!_{xx}()$, alors l'équation 7.6 devient

$$\mathbf{w}(\delta, i)^*\mathbf{w}(\delta, j) = \begin{cases} 0 & \text{si } i \neq j \\ 1 & \text{sinon} \end{cases}$$

dès que $!_l(1) = 1$, et l'équation 7.5 se déduit de

$$\mathbf{w}(\delta, 1)\mathbf{w}(\delta, 1)^* + \mathbf{w}(\delta, 2)\mathbf{w}(\delta, 2)^* = 1$$

effacement Soient c une s -boîte réduite sur une cellule d'effacement, et φ un chemin pax/pax dans c . On devrait avoir $!_{xx}(\mathbf{w}(\varphi)) = 0$. Or, le même chemin présent dans une autre s -boîte n'étant pas effacée devrait vérifier $!_{xx}(\mathbf{w}(\varphi)) \neq 0$ car c'est un chemin du réseau normal. On ne peut donc pas tenir compte de l'effacement pour obtenir une pondération fortement fidèle ou même tout simplement fidèle. L'approche la plus raisonnable serait de ne pas annuler le poids de tels chemins, quitte à procéder à une étape supplémentaire a posteriori afin d'éliminer les poids en trop.

enfouissement Soient b et b' deux s -boîtes, où b' s'enfouit dans b , on doit avoir

$$!_{lx}(\mathbf{w}(\varphi)) * !_{ll}(\mathbf{w}(\varphi')) !_{lx}(\mathbf{w}(\varphi'')) = !_{ll}(\mathbf{w}(\varphi) * !_{ll}(\mathbf{w}(\varphi')) \mathbf{w}(\varphi'')) \quad (7.8)$$

$$!_{lx}(\mathbf{w}(\varphi)) !_{lx}(\mathbf{w}(\varphi')) = !_{lx}(!_{lx}(\mathbf{w}(\varphi)) \mathbf{w}(\varphi')) \quad (7.9)$$

$$!_{xx}(\mathbf{w}(\varphi)) = !_{xx}(!_{xx}(\mathbf{w}(\varphi))) \quad (7.10)$$

L'équation 7.10 est exprimée pour φ chemin pax/pax dans b' . L'équation 7.9 est conséquence de

$$!_{lx}(x) = !_{lx}(!_{lx}(x))$$

agrégation Soient b et b' deux s -boîtes, avec b' s'agrégeant dans b . On doit avoir

$$!_{lx}(\mathbf{w}(\varphi)) * !_{ll}(\mathbf{w}(\varphi')) !_{lx}(\mathbf{w}(\varphi'')) = !_{ll}(\mathbf{w}(\varphi) * \mathbf{w}(\varphi') \mathbf{w}(\varphi'')) \quad (7.11)$$

$$!_{lx}(\mathbf{w}(\varphi)) !_{lx}(\mathbf{w}(\varphi')) = !_{lx}(\mathbf{w}(\varphi) \mathbf{w}(\varphi')) \quad (7.12)$$

$$!_{xx}(\mathbf{w}(\varphi)) = !_{xx}(\mathbf{w}(\varphi)) \quad (7.13)$$

Notons que l'équation 7.13 est une tautologie, et que l'équation 7.12 est conséquence du fait que $!_{lx}()$ soit un morphisme. Seule l'équation 7.11 est à vérifier. Une manière simple de l'obtenir est de poser $!_{lx}(x) = !_{ll}(x)$.

Remarque 7.10 *Les principaux freins à l'obtention d'une pondération fortement fidèle sont tous liés aux chemins pax/pax ou pal/pax.*

7.2 Géométrie de l'Interaction

Définition 7.11 *Soit $\mathcal{B}$ une bibliothèque close de réseaux d'interaction, on appelle Géométrie de l'Interaction, GdI en abrégé, de $\mathcal{B}$ une pondération fortement fidèle pour $\mathcal{B}$.*

Voilà les grandes étapes de la réalisation d'une GdI :

1. donner une présentation par générateurs et relations d'un monoïde D .
2. prouver que D est normal. On se ramène en général à montrer que D a un modèle non-trivial et à déduire la normalité de la forme des relations. Dans les cas usuels on s'arrange pour que D soit en fait un monoïde simplifiable.
3. définir une pondération générique et simple pour $\mathcal{B}$ à valeur dans D .
4. vérifier l'équation 7.1 pour chaque règle.

Définition-Propriété 7.12 *Soient $\mathbf{w}$ une GdI d'une bibliothèque $\mathcal{B}$ complète vis-à-vis d'un typage, et $R \in \mathfrak{R}(\mathcal{S}')^t$ un réseau bien typé associé. On suppose que $\mathbf{w}$ est à valeur dans $\mathbb{C}[\mathbf{M}]$.*

On appelle GdI de R , ou encore résultat de l'exécution de R , l'invariant de la réduction $\rightarrow_{\mathcal{B}}$

$$\text{Ex}(R) = \sum_{\varphi \in \mathfrak{P}_f(R)} \mathbf{w}(\varphi) = \sum_{\varphi \in \mathfrak{P}_{\mathbf{r}_f}(R)} \mathbf{w}(\varphi) \in \mathbb{C}[\mathbf{M}]$$

Preuve Supposons que $R \rightarrow_{\mathcal{R}} R'$, on a $\text{Ex}(R) = \sum_{\varphi \in \mathfrak{P}_f(R)} \mathbf{w}(\varphi) = \sum_{\varphi \in \mathfrak{P}_f(R)} \mathbf{w}(\delta_{\mathcal{R}}(\varphi))$ car $\mathbf{w}$ est fortement fidèle. Or, en vertu du fait 6.8 on a ainsi tous les chemins de $\mathfrak{P}_f(R')$ et $\text{Ex}(R) = \text{Ex}(R')$.

Maintenant supposons que R_0 soit la forme normale de R . On a $\text{Ex}(R) = \text{Ex}(R_0)$. Mais $\mathfrak{P}_f(R_0)$ étant fini on en déduit que $\text{Ex}(R)$ ne compte qu'un nombre fini de termes dans $\mathbf{M}$, et appartient donc à l'algèbre $\mathbb{C}[\mathbf{M}]$. $\blacktriangleright$

Remarque 7.13 Remarquons que si $\varphi \in \mathfrak{P}_f(R)$ alors $\varphi^r \in \mathfrak{P}_f(R)$, et ainsi $\text{Ex}(R)^{\star} = \text{Ex}(R)$. On peut donc se contenter de calculer $\text{Ex}_{\frac{1}{2}}(R) = \sum_{\varphi \in P} \mathbf{w}(\varphi)$ où $P \subset \mathfrak{P}_f(R)$ ne contient qu'un seul des chemins à inversion près. On récupère $\text{Ex}(R)$ par le calcul $\text{Ex}(R) = \text{Ex}_{\frac{1}{2}}(R) + \text{Ex}_{\frac{1}{2}}(R)^{\star}$.

7.3 Exemples de GdI

7.3.1 Notations synthétiques pour les relations usuelles

Dans la suite nous utiliserons les notations suivantes :

$$\begin{aligned} p \perp q &\iff p^{\star}q = q^{\star}p = 0 \\ p \overline{\perp} q &\iff p \perp q, p^{\star}p = q^{\star}q = 1 \\ p \overline{\perp}^+ q &\iff p^{\star}p = q^{\star}q = 1, pp^{\star} + qq^{\star} = 1 \\ u \leftrightarrow v &\iff uv = vu \\ E \leftrightarrow F &\iff \forall u \in E, \forall v \in F, u \leftrightarrow v \\ u \overset{\star}{\leftrightarrow} v &\iff \{u\} \leftrightarrow \{v, v^{\star}\} \\ E \overset{\star}{\leftrightarrow} F &\iff \forall u \in E, \forall v \in F, u \overset{\star}{\leftrightarrow} v \end{aligned}$$

Si $p \perp q$ (resp. $p \overline{\perp} q$) on dira que p et q sont orthogonaux (resp. pleinement orthogonaux). Les propriétés $\leftrightarrow$ et $\overset{\star}{\leftrightarrow}$ sont appelées propriétés de commutation.

La relation $p \overline{\perp}^+ q$ est donc une relation quotientant l'algèbre du monoïde construit. Notons que dans ce quotient $p \overline{\perp}^+ q \Rightarrow p \overline{\perp} q$.

7.3.2 GdI de MLL

La GdI de MLL, définie en 4.3, est une des plus simples possibles qui soit non triviale. Soit

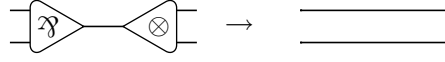
$$\mathbf{M} = \langle p, q \mid p \overline{\perp} q \rangle_0^1$$

il est normal et on peut montrer qu'il vérifie la propriété $ab^{\star}$. On pose

$$\mathbf{w}(\mathfrak{A}, 1) = p \quad \mathbf{w}(\mathfrak{A}, 2) = q$$

$$\mathbf{w}(\otimes, 1) = q \quad \mathbf{w}(\otimes, 2) = p$$

Rappelons la règle de réduction :



Les équations que doit vérifier w pour être fortement fidèle sont donc

$$\begin{aligned} w(c_{1,1}c_{2,2}^r) &= p^*p = 1 & w(c_{1,2}c_{2,1}^r) &= q^*q = 1 \\ w(c_{1,1}c_{2,1}^r) &= q^*p = 0 \text{ et } w(c_{1,2}c_{2,2}^r) &= q^*p = 0 \end{aligned}$$

Or, ce sont exactement les relations qui ont servi à construire M .

7.3.3 GdI de ELL

Soit

$$M = \left\langle p, q, r, s \mid \begin{array}{l} p \perp q, r \perp^+ s \\ !(x) \leftrightarrow \{r, s\} \end{array} \right\rangle_0^1$$

Notons que la relation $rr^* + ss^* = 1$ n'est pas directement une relation de monoïde. En effet, on ne peut quotienter selon cette relation que dès lors qu'on a introduit la somme en prenant l'algèbre du monoïde. Cependant, cette relation entraînant $r^*s = 0$ il est naturel de la considérer dans la présentation de M .

M est un monoïde normal et vérifiant la propriété ab^* .

On pose

$$\begin{aligned} w(\mathcal{A}, 1) &= p & w(\mathcal{A}, 2) &= q \\ w(\otimes, 1) &= q & w(\otimes, 2) &= p \\ !_l(x) &= !(x) & !_{lx}(x) &= !(x) & !_{xx}(x) &= !(x) \\ w(?^2, 1) &= r & w(?^2, 2) &= s \end{aligned}$$

On vérifie rapidement que la pondération ainsi construite est fortement fidèle.

7.3.4 GdI de MELL

$$M = \left\langle p, q, r, s, d, t \mid \begin{array}{l} p \perp q, r \perp s \\ d^*d = t^*t = 1, \\ !(x) \leftrightarrow \{r, s\}, \\ !(x)t = t!(x), !(x)d = dx \end{array} \right\rangle_0^1$$

On pose

$$\begin{aligned} w(\mathcal{A}, 1) &= p & w(\mathcal{A}, 2) &= q \\ w(\otimes, 1) &= q & w(\otimes, 2) &= p \\ !_l(x) &= !(x) & !_{lx}(x) &= t!(x) & !_{xx}(x) &= t!(x)t^* \\ w(?^2, 1) &= r & w(?^2, 2) &= s & w(?^1, 1) &= d \end{aligned}$$

w ainsi défini n'est pas fortement fidèle en vertu de la remarque 7.10. Cependant, on peut supprimer tous les problèmes liés aux boîtes si l'on ne réduit que des boîtes sans portes auxiliaires.

Théorème 7.14 *Soit $\mathfrak{R}_?$ l'ensemble des réseaux de MELL dont le type ne comprend pas de ?.*

La réduction de MELL est une réduction sur $\mathfrak{R}_?$.

Si $r \in \mathfrak{R}_?$ contient au moins un fil actif connecté au port principal d'une !-boîte, une promotion, alors r contient un fil actif connecté au port principal d'une promotion n'ayant pas de portes auxiliaires.

w est une pondération fortement fidèle sur $\mathfrak{R}_?$.

La preuve de ce fait se trouve dans l'article fondateur [Gir89a] de Jean-Yves Girard.

7.3.5 GdI de MELL avec connecteurs synthétiques

Afin de donner une pondération dans ce nouveau cadre, nous allons abandonner pendant un temps la genericité. En effet, la pondération d'une cellule va maintenant dépendre de la portance de ses ports auxiliaires dans le cas des cellules ?. Notons $h(c)$ la hauteur de la cellule c , et $l_i(c)$ la portance de son i -ème auxiliaire si c est une cellule ?.

Soit R un réseau de SMELL, si on numérote $e_1, \dots, e_n$ ses cellules ?, on considère un ensemble de générateurs $((\beta_{i,j})_{j \in \{1, \dots, \alpha(e_i)\}})_{i \in \{1, \dots, n\}}$, appelé *paquet de constantes exponentielles*, où $\alpha(e_i)$ est le nombre des ports auxiliaires de e_i . On note $c(\beta_{i,j}) = l_j(e_i)$ qu'on appelle *charge* de $\beta_{i,j}$.

Soit

$$\mathbf{M}_R = \left\langle p, q, \beta_{i,j} \mid \begin{array}{l} p \bar{\perp} q, \\ \forall j \neq k, \beta_{i,j} \bar{\perp} \beta_{i,k}, \\ !(x)\beta_{i,j} = \beta_{i,j}!^{c(\beta_{i,j})}(x) \end{array} \right\rangle_0^1$$

Ce monoïde est simplifiable et a la propriété ab^* .

On pose $w(e_{i,j}) = \beta_{i,j}$ ainsi que $!_l(x) = !_{lx}(x) = !_{xx}(x) = !(x)$.

Si $R \rightarrow R'$ par une réduction, Regnier et Danos ont montré dans [DR95] qu'il existe un modèle non trivial $\mathbf{M}$ de $\mathbf{M}_R$ et $\mathbf{M}_{R'}$ dans lequel $\mathbf{M} \models w(\varphi) = w(\delta(\varphi))$. C'est-à-dire un résultat de persistance forte vis-à-vis de M . Or M étant non trivial on peut appliquer le fait 5.16, et en déduire que M est fidèle. Donc, ce résultat de persistance forte relative devient un résultat général de persistance.

Remarque 7.15 *Il est légitime de vouloir comparer ces deux GdI de MELL. La première nous donne un ensemble, d'apparence restreint, de réseaux pour lesquels on a préservation algébrique du poids lors de la réduction. La seconde nous permet de caractériser les chemins persistants partout.*

Chapitre 8

λ -calcul

$$\begin{array}{c}
 \vdots \\
 ((\lambda x.xx), \{\}) \\
 \hline
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Push} \\
 \hline
 (xx, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\})\}) \\
 \star \epsilon \quad \text{Pop} \\
 \hline
 ((\lambda x.xx), \{\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\}) \bullet \epsilon \quad \text{Push} \\
 \hline
 (xx, \{x \mapsto (x, \{x \mapsto (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\}), \{\}), \{\}), \{\})\}) \\
 \star \epsilon \quad \text{Pop} \\
 \hline
 ((\lambda x.xx), \{\}) \\
 \star(x, \{x \mapsto ((\lambda x.xx), \{\}), \{\})\}) \bullet \epsilon \quad \text{Jump} \\
 \hline
 (x, \{x \mapsto ((\lambda x.xx), \{\}), \{\})\}) \\
 \star(x, \{x \mapsto ((\lambda x.xx), \{\}), \{\})\}) \bullet \epsilon \quad \text{Push} \\
 \hline
 (xx, \{x \mapsto ((\lambda x.xx), \{\}), \{\})\}) \\
 \star \epsilon \quad \text{Pop} \\
 \hline
 ((\lambda x.xx), \{\}) \\
 \star((\lambda x.xx), \{\}) \bullet \epsilon \quad \text{Push} \\
 \hline
 ((\lambda x.xx)(\lambda x.xx), \{\}) \\
 \star \epsilon
 \end{array}$$

Le λ-calcul peut être vu comme un langage d'assembleur pour la programmation fonctionnelle : inutilisable pour écrire directement des programmes mais fondamental pour leur étude. En effet, tous les fonctions calculables s'encodant dans le λ-calcul, il suffit de savoir évaluer correctement le λ-calcul pour savoir presque tout évaluer.

8.1 Définitions

On se donne dans ce chapitre un ensemble dénombrable de variables $x, y, \dots$. On appelle λ-termes les termes définis par la grammaire

$$u, v ::= uv \mid \lambda x. u \mid x$$

L'ensemble des λ-termes est noté Λ .

On appelle contextes les termes définis par

$$u, v ::= uv \mid \lambda x. u \mid x \mid \langle \rangle$$

et contenant exactement une occurrence du symbole $\langle \rangle$. L'ensemble des contextes est noté $\Lambda\langle \rangle$. Étant donné un contexte C , on note $C\langle t \rangle$ le λ-terme obtenu en remplaçant $\langle \rangle$ par t . Une relation binaire $\mathcal{R}$ sur Λ est dite *passant au contexte* lorsque

$$\forall C \in \Lambda\langle \rangle, \forall u, v \in \Lambda, u \mathcal{R} v \Rightarrow C\langle u \rangle \mathcal{R} C\langle v \rangle$$

On appelle β-réduction, et l'on note $\rightarrow_\beta$, la plus petite relation passant au contexte telle que

$$(\lambda x. t)u \rightarrow_\beta t[u/x] \tag{8.1}$$

où $x \notin u$ et $\bullet[\bullet/\bullet]$ signifie

$$(uv)[M/x] = u[M/x]v[M/x]$$

$$(\lambda y. t)[M/x] = \begin{cases} \lambda y. t & \text{si } x = y \\ \lambda y. t[M/x] & \text{sinon} \end{cases} \quad y[M/x] = \begin{cases} M & \text{si } x = y \\ y & \text{sinon} \end{cases}$$

Le motif à gauche de (8.1) est appelé un *redex*. Un terme sans redex est dit *en forme normale*.

Étant donné un terme possédant plusieurs redexs, une procédure déterminant de manière systématique lequel réduire est appelé *stratégie de réduction*. La stratégie la plus simple consiste à réduire le redex le plus à gauche, parmi les innombrables autres familles de stratégies citons l'appel-par-nom et l'appel-par-valeur, caractéristiques du fait qu'à une famille de stratégie peut correspondre un paradigme complet.

8.2 Exemples

On donne ici quelques exemples de λ-termes ainsi que leur réductions. On notera $\lambda x_1 \dots x_n. t$ le terme $\lambda x_1. \dots \lambda x_n. t$.

Le λ-terme identité $\lambda x. x$ dont la réduction suivante justifie l'appellation :

$$(\lambda x. x)t \rightarrow_\beta t$$

On pose $\delta = \lambda x. xx$ et $\Omega = \delta\delta$. On a $\Omega \rightarrow_\beta \Omega$.

Pour tout $n \in \mathbb{N}$, on pose $\bar{n} = \lambda f x. f^n x$ où $f^0 x = x$ et $f^{n+1} x = f(f^n x)$, ces termes sont appelés entiers de Church. On a, avec $m, n \neq 0$ dans la première réduction,

$$\bar{n} \bar{m} \rightarrow_{\beta}^* \overline{m^n}, \quad (\lambda n m f x. n f(m f x)) \bar{n} \bar{m} \rightarrow_{\beta}^* \overline{n + m}, \quad (\lambda n m f x. n(m f x)) \bar{n} \bar{m} \rightarrow_{\beta}^* \overline{nm}$$

8.3 La KAM

Une stratégie de réduction ne donne pas de manière concrète d'évaluer les termes : il est toujours nécessaire de recourir de manière extérieure à un mécanisme de substitution. On appelle machine d'évaluation un procédé réalisant entièrement une stratégie de réduction. La machine la plus simple est sans doute la KAM, pour Krivine Abstract Machine, servant de base dans les travaux de Krivine sur la réalisabilité classique.

La KAM n'évalue pas directement des termes mais des *clôtures*, c'est-à-dire des paires (t, E) où t est un terme et E est une fonction partielle de domaine fini des variables dans les clôtures, appelée *environnement*. Un état de la KAM est une paire $c \star \pi$ où c est une clôture et π une pile de clôtures, et un état initial pour évaluer le terme t est $(t, E) \star \epsilon$, où E a pour domaine les variables libres de t et ϵ est la constante de fond de pile. On note $\bullet$ l'opérateur de concaténation de la pile.

L'exécution de la KAM s'effectue selon les règles suivantes, orientées du bas vers le haut :

$$\frac{(t, E \cup \{x \mapsto c\}) \star \pi}{(\lambda x. t, E) \star c \bullet \pi} \text{ Pop } (x \text{ non libre dans } E) \quad \frac{(u, E) \star (v, E) \bullet \pi}{(u \ v, E) \star \pi} \text{ Push}$$

$$\frac{(t, E') \star \pi}{(x, E \cup \{x \mapsto (t, E')\}) \star \pi} \text{ Jump}$$

Ces trois règles sont suffisantes pour étudier l'évaluation mais pour pouvoir vraiment calculer des réduits avec cette machine il est nécessaire de faire quelques modifications.

La contrainte sur la règle Pop signifie que x ne doit pas apparaître libre dans un terme d'une clôture image de E , cette condition est quelquefois appelée *fraîcheur*. Elle se contourne facilement en rajoutant une étape explicite d' α -conversion :

$$\frac{(\lambda y. t[y/x], E) \star \pi}{(\lambda x. t, E) \star \pi} \alpha\text{-conv } (x \text{ libre dans } E, y \text{ non libre dans } E \text{ et } t)$$

Pour s'appliquer, la règle Pop nécessite que la pile contienne au moins une clôture. Cette restriction revient à ne considérer que des réductions vers des variables libres. Une pile vide correspond donc à un λ en tête, et on peut continuer le calcul en introduisant une règle :

$$\frac{(t, E) \star \epsilon}{(\lambda x. t, E) \star \epsilon} \text{ Head } \lambda x$$

Clore l'environnement de départ permet d'assurer que la règle Jump s'applique toujours, cela signifie également que le calcul ne s'arrête jamais. Finir sur une variable x n'apparaissant pas dans l'environnement correspond à produire une valeur de retour qu'on relit en accumulant les λ produit par la règle précédente et en évaluant le reste de la pile. De la manière d'évaluer ces termes restants dépend le type de forme normale que l'on va obtenir. Essentiellement, deux choix se présentent : soit complètement les évaluer à l'aide de nouvelles instances de la

KAM, soit simplement effectuer les substitutions des variables liés par l'environnement. Dans le premier cas on obtiendra la forme normale complète et dans l'autre la forme normale de tête.

Pour calculer la forme normale de tête il suffit d'ajouter la règle suivante :

$$\frac{\lambda x_1 \dots \lambda x_n. x \ u_1[e_1] \dots u_m[e_m]}{(x, \{\}) \star (u_1, e_1) \bullet \dots \bullet (u_m, e_m) \bullet \epsilon} \text{ Head Return}$$

$$\frac{\vdots}{\text{Head } \lambda x_n}$$

$$\frac{\vdots}{\text{Head } \lambda x_1}$$

$$\vdots$$

où l'on ne fait qu'effectuer des substitutions de contexte dans les u_i .

Pour obtenir la forme normale il faut à la place effectuer inductivement la réduction des u_i . La règle à ajouter pour le calcul de la forme normale est alors :

$$\frac{\lambda x_1 \dots \lambda x_n. x \ u_1 \dots u_m \text{ Return} \quad \star \quad \frac{u_1}{c_1} \text{ Return} \quad \bullet \dots \bullet \quad \frac{u_m}{c_m} \text{ Return} \quad \bullet \epsilon}{\vdots} \text{ Subcall}$$

$$\frac{\vdots}{\text{Head } \lambda x_n}$$

$$\frac{\vdots}{\text{Head } \lambda x_1}$$

$$\vdots$$

où cela signifie que u_i est le résultat final, par une règle Return de l'évaluation de c_i .

La notation $t[e]$ signifie $t[t'_1/x_1] \dots [t'_n/x_n]$ pour $\text{dom}(e) = \{x_1, \dots, x_n\}$, $e(x_i) = (t_i, e_i)$ et $t'_i = t_i[e_i]$. Notons que cette opération de substitution est bien définie pour tous les termes, contrairement à la règle précédente qui nécessite que chaque argument soit normalisable à gauche. Ainsi, pour le terme $x\Omega$, où $\Omega = (\lambda x.xx)(\lambda x.xx)$, Return essaye, sans succès¹, d'évaluer Ω , alors que le terme est renvoyé directement par Head Return. Les figures Fig. 8.1 et Fig. 8.2 présentent des évaluations du terme $\bar{2} \bar{2}$ par la KAM muni de la règle Return et de la règle Head Return.

8.4 Traduction du λ-calcul dans SMELL

Nous donnons ici la traduction des termes dans les réseaux SMELL, voir section 4.4.2, comme elle apparaît dans la thèse [Reg92] de Laurent Regnier.

On note $\text{FV}(t)$ le *multi-ensemble* des occurrences de variables libres dans le terme t . On note $\text{occ}_x(t)$ le multi-ensemble des occurrences libres de x dans t , en sorte que $\text{FV}(t) = \sum_x \text{occ}_x(t)$.

Les réseaux obtenus par traduction seront munis d'un étiquetage de leur ports libres par des variables ou par la constante o .

Soit t un λ-terme, on définit la traduction $t^\bullet$ par induction :

variable si $t = x$ variable, on pose

$$x^\bullet = \text{arc from } x \text{ to } o$$

¹ cf. la page de garde du chapitre.

$$\begin{array}{c}
\frac{\frac{\frac{x_0}{(x_0, \{\})} \text{Return}}{(x, \{\})} \text{Return} \quad \frac{\frac{x_0}{(x_0, \{\})} \text{Return}}{(x_0, \{x_0 \mapsto (x_0, \{\}), \{\})} \text{Jump}}{\star \epsilon} \text{Subcall} \\
\frac{(x, \{\})}{\star(x_0, \{x_0 \mapsto (x_0, \{\}), \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (x, \{\}), \{\})}{\star(x_0, \{x_0 \mapsto (x_0, \{\}), \{\}) \bullet \epsilon} \text{Push} \\
\frac{\frac{x^2 x_0}{(x, \{\})} \text{Return} \quad \frac{(f x_0, \{x_0 \mapsto (x_0, \{\}), f \mapsto (x, \{\}), \{\})}{\star \epsilon} \text{Subcall}}{\star(f x_0, \{x_0 \mapsto (x_0, \{\}), f \mapsto (x, \{\}), \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (x, \{\}), \{\})}{\star(f x_0, \{x_0 \mapsto (x_0, \{\}), f \mapsto (x, \{\}), \{\}) \bullet \epsilon} \text{Push} \\
\frac{(f^2 x_0, \{x_0 \mapsto (x_0, \{\}), f \mapsto (x, \{\}), \{\})}{\star \epsilon} \text{Pop} \\
\frac{(\lambda x_0. f^2 x_0, \{f \mapsto (x, \{\}), \{\})}{\star(x_0, \{\}) \bullet \epsilon} \alpha\text{-conv } x \rightarrow x_0 \\
\frac{(\lambda x. f^2 x, \{f \mapsto (x, \{\}), \{\})}{\star(x_0, \{\}) \bullet \epsilon} \text{Pop} \\
\frac{(\lambda f x. f^2 x, \{\})}{\star(x, \{\}) \bullet (x_0, \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\})}{\star(x, \{\}) \bullet (x_0, \{\}) \bullet \epsilon} \text{Push} \\
\frac{(f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\})}{\star(x_0, \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\})}{\star(x_0, \{\}) \bullet \epsilon} \text{Push} \\
\frac{(f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\})}{\star \epsilon} \text{Jump} \\
\frac{\frac{x^3 x_0}{(x, \{\})} \text{Return} \quad \frac{(x_0, \{x_0 \mapsto (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}), \{\})}{\star \epsilon} \text{Subcall}}{\star(x_0, \{x_0 \mapsto (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}), \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (x, \{\}), \{\})}{\star(x_0, \{x_0 \mapsto (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}), \{\}) \bullet \epsilon} \text{Push} \\
\frac{\lambda x x_0. x^4 x_0}{(x, \{\})} \text{Return} \quad \frac{(f x_0, \{x_0 \mapsto (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}), \{\})}{\star \epsilon} \text{Subcall}}{\star(f x_0, \{x_0 \mapsto (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}), \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (x, \{\}), \{\})}{\star(f x_0, \{x_0 \mapsto (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}), \{\}) \bullet \epsilon} \text{Push} \\
\frac{(f^2 x_0, \{x_0 \mapsto (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}), \{\})}{\star \epsilon} \text{Pop} \\
\frac{(\lambda x_0. f^2 x_0, \{f \mapsto (x, \{\}), \{\})}{\star(f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}) \bullet \epsilon} \alpha\text{-conv } x \rightarrow x_0 \\
\frac{(\lambda x. f^2 x, \{f \mapsto (x, \{\}), \{\})}{\star(f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}) \bullet \epsilon} \text{Pop} \\
\frac{(\lambda f x. f^2 x, \{\})}{\star(x, \{\}) \bullet (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\})}{\star(x, \{\}) \bullet (f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}) \bullet \epsilon} \text{Push} \\
\frac{(f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\})}{\star(f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\})}{\star(f x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\}) \bullet \epsilon} \text{Push} \\
\frac{(f^2 x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\})}{\star \epsilon} \text{Head } \lambda x_0 \\
\frac{(\lambda x_0. f^2 x_0, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\})}{\star \epsilon} \alpha\text{-conv } x \rightarrow x_0 \\
\frac{(\lambda x. f^2 x, \{f \mapsto (f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}), \{\})}{\star \epsilon} \text{Pop} \\
\frac{(\lambda f x. f^2 x, \{\})}{\star(f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}) \bullet \epsilon} \text{Jump} \\
\frac{(f, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\})}{\star(f x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\}) \bullet \epsilon} \text{Push} \\
\frac{(f^2 x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\})}{\star \epsilon} \text{Head } \lambda x \\
\frac{(\lambda x. f^2 x, \{f \mapsto (\lambda f x. f^2 x, \{\}), \{\})}{\star \epsilon} \text{Pop} \\
\frac{(\lambda f x. f^2 x, \{\})}{\star(\lambda f x. f^2 x, \{\}) \bullet \epsilon} \text{Push} \\
\frac{((\lambda f x. f^2 x)(\lambda f x. f^2 x), \{\})}{\star \epsilon}
\end{array}$$

Figure 8.1: Evaluation de $(\lambda f x. f^2 x)(\lambda f x. f^2 x)$ par la KAM + Return

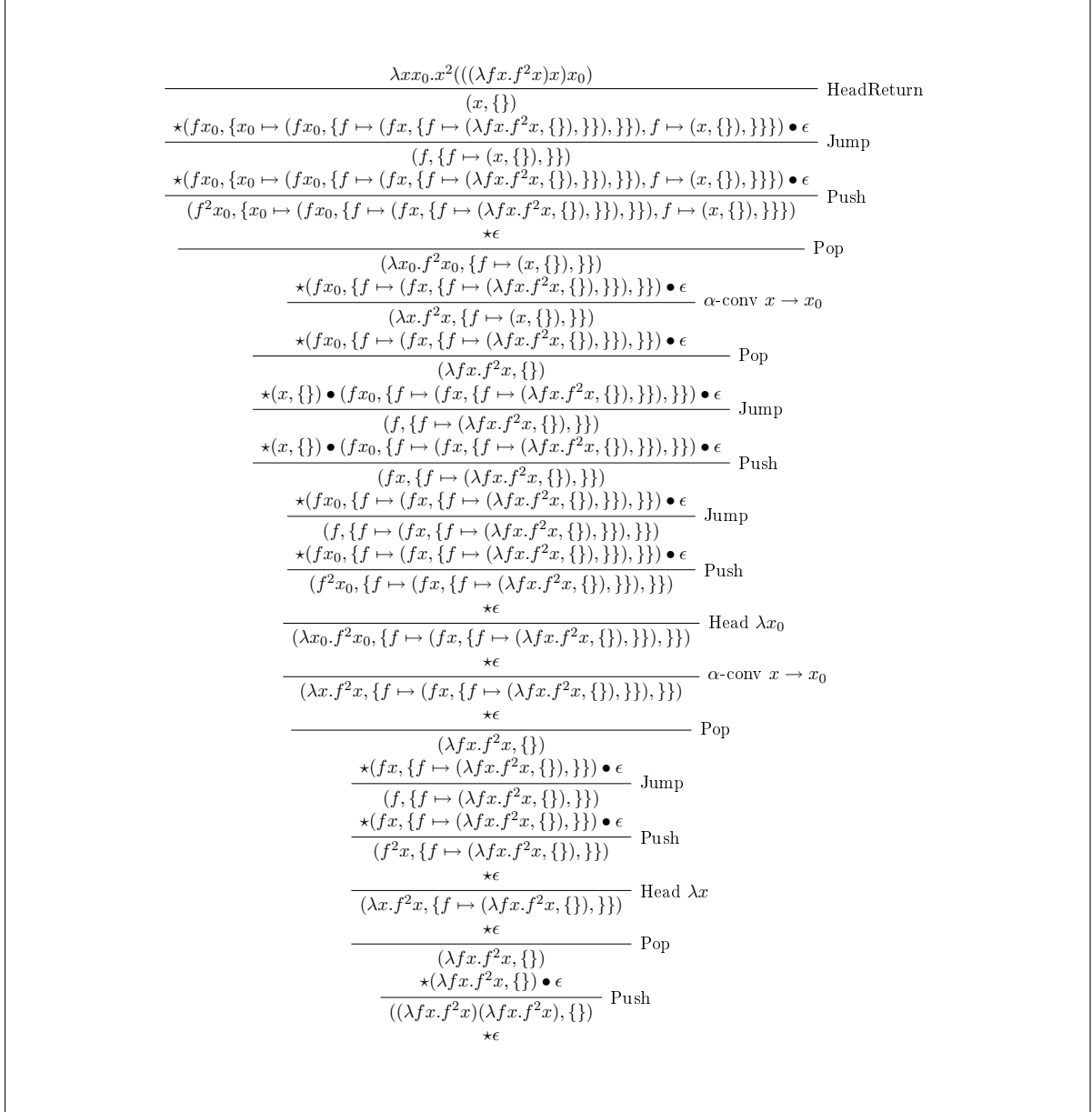
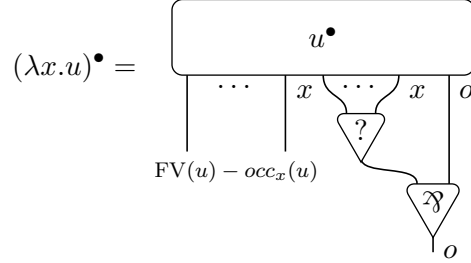
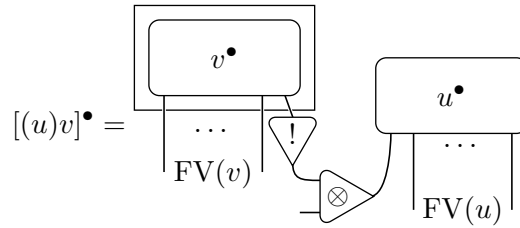


Figure 8.2: Evaluation de $(\lambda f x. f^2 x)(\lambda f x. f^2 x)$ par la KAM + Head Return

abstraction si $t = \lambda x.u$, on pose



application si $t = (u)v$, on pose



Si $t \rightarrow_\beta t'$ on a $t^\bullet \rightarrow^2 t'^\bullet$, l'inverse n'étant pas vrai dans la mesure où la réduction n'est plus atomique dans les réseaux. On aurait pu par exemple partir de $t^\bullet$ et réduire deux redex $\wp/\otimes$ obtenant ainsi un réseau qui n'est pas la traduction d'un λ -terme.

8.4.1 Typage pur

Les traductions de termes du λ -calcul pur ne sont pas tous typables au sens usuel. On introduit une nouvelle discipline de typage sur deux atomes O et I ainsi que l'équation $O = ?I \wp O$. La traduction d'un λ -terme est alors bien typé, le port libre d'étiquette o a le type O et chaque port étiqueté par une variable a le type I .

8.4.2 GdI

Notons qu'avec le typage pur il n'est plus possible de récupérer directement une condition précise de préservation de la GdI comme celle du théorème 7.14. Remarquons que si t a pour forme normale une variable, alors $\text{Ex}(t^\bullet)$ est un invariant.

8.4.3 Exemples

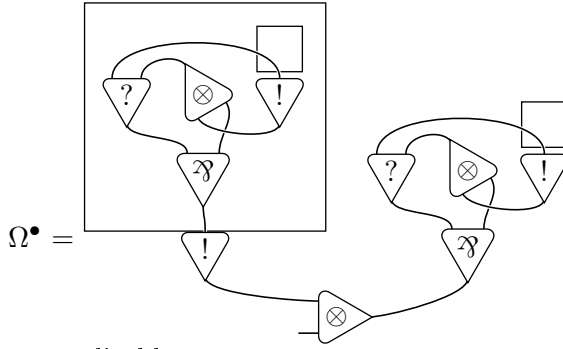
On ne donne ici que la moitié des exécutions en vertu de la remarque 7.13.

$$(\lambda x.x)^\bullet =$$

et $\text{Ex}_{\frac{1}{2}}((\lambda x.x)^\bullet) = p\alpha^*q^*$ où α est de charge 0.

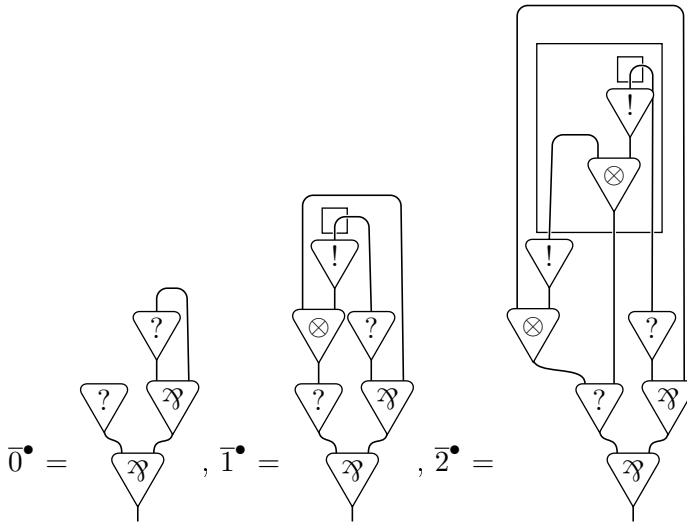
$$\delta^\bullet =$$

et $\text{Ex}_{\frac{1}{2}}(\delta^\bullet) = pz_0qq^* + pz_1p^*z_0^*p^*$ où z_i est de charge i .



dont on ne peut calculer la GdI car il n'est

pas normalisable.



On voit apparaître une récurrence évidente dans l'expression de ces réseaux nous permettant d'obtenir $\bar{n}^\bullet$. Du point de vue de la GdI, supposons donné deux ensembles de constantes exponentielles $(f_i)_{i \in \mathbb{N}}$ et $(x_i)_{i \in \mathbb{N}}$ telles que la charge de f_i et x_i soit i . On a

$$\text{Ex}_{\frac{1}{2}}(\bar{n}^\bullet) = pf_0qq^*q^* + \sum_{k=1}^{n-1} pf_k!^{k-1}(! (q)p^*)f_{k-1}^*p^* + qp x_n!^{n-1}(p^*)f_{n-1}^*p^*$$

Une grande partie du chapitre suivant sera consacré à l'étude de la GdI de ces termes, et plus particulièrement de leurs sens.

Chapitre 9

Concision de la GdI des λ -termes

Dans tout ce chapitre on étudie SMELL est plus particulièrement les traductions de λ -termes. On notera les chemins avec la notion traditionnelle des vecteurs $\overrightarrow{\phi}$ pour les distinguer des poids. On notera $\overleftarrow{\phi} = \overrightarrow{\phi}^r$.

9.1 Concision

Définition 9.1 Soit R un réseau dont la GdI est $\sum_{i=1}^n w_i$ avec les $w_i \in \mathbb{M}$. On appelle concision de la GdI de R l'entier $|R| = \frac{n}{2}$.

Notons que la définition est valide car n est pair en vertu de la remarque 7.13. Pour un λ -terme t on note abusivement $|t| = |t^\bullet|$.

Si $t \rightarrow_\beta^* t'$ nous n'avons pas nécessairement $\text{Ex}(t^\bullet) = \text{Ex}(t'^\bullet)$. Ce qui est encore plus étonnant est que $|t|$ n'est pas non plus nécessairement égal à $|t'|$. Cela est surprenant car, s'il existe $u_1, \dots, u_n$ tel que $tu_1 \dots u_n \rightarrow_\beta^* x$ où x est une variable, alors $\text{Ex}((tu_1 \dots u_n)^\bullet) = \text{Ex}((t'u_1 \dots u_n)^\bullet)$. Cela signifie qu'il y a dans la GdI de t autant d'information que dans la GdI de t' . On voit donc pourquoi $|t|$ a été appelé concision : plus ce nombre est petit plus la GdI a *compressé* l'information.

Le problème qui se pose alors naturellement est le suivant :

Soit t_0 un λ -terme en forme normale, peut on calculer $\inf_{t \rightarrow_\beta^* t_0} |t|$?
Cette borne inférieure est-elle atteinte ?

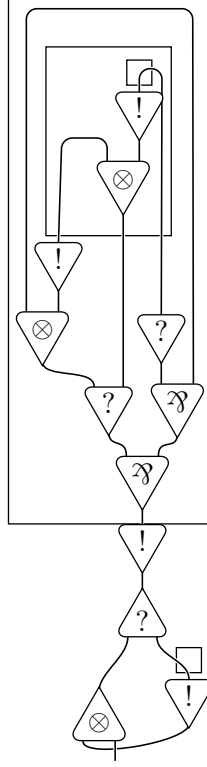
Ces questions sont très complexes et pour l'instant sans solution. Dans ce chapitre on va essayer de mesurer la difficulté du problème sur un cas caractéristique : les entiers de Church. Les raisons de ce choix sont multiples, d'une part la GdI de ces termes a une signification opérationnelle très claire comme nous le verront, et d'autre part on peut y exprimer un calcul simple qui fait apparaître une *compression* spectaculaire : l'exponentiation. De plus, formant une famille de concision croissante, $|\overline{n}| = n + 1$, ils permettent l'expression de propriétés asymptotiques, premier angle d'attaque de notre problème.

9.2 Un calcul de $|\delta \overline{n}|$

Avant de commencer, admettons qu'aucun phénomène de compression n'est dû à l'utilisation de δ , et on a $|\delta \overline{n}| = |\overline{n} \ \overline{n}|$.

La raison de l'utilisation de δ est qu'il nous permet de n'utiliser qu'un jeu de poids pour les deux $\bar{n}$ et cela avec un coût minime dans la mesure où la GdI de δ est très simple.

On va de plus faire une première réduction multiplicative avant de calculer la GdI. Ainsi, le réseau dont on calculera la sémantique pour $\delta\bar{2}$ sera



9.2.1 Chemins et poids dans $\bar{n}$

Rappelons que le calcul de $\text{Ex}(\bar{n})$ a été effectué à la fin du chapitre précédent. Nous allons ici décomposer cette GdI. On note $\vec{f_0}, \dots, \vec{f_{n-1}}, \vec{x_n}$ les chemins de $\mathfrak{P}_f(R)$, où R est le réseau contenant une unique !-boîte de contenu $\bar{n}^\bullet$, tels que

$$\mathbf{w}(\vec{f_0}) = !(pf_0qq^*q^*)$$

$$\mathbf{w}(\vec{f_k}) = !(pf_k!^{k-1}(! (q)p^*)f_{k-1}^*p^*) \text{ pour } n > k > 0$$

$$\mathbf{w}(\vec{x_n}) = !(qp x_n!^{n-1}(p^*)f_{n-1}^*p^*)$$

Remarquons que seul le poids de $\vec{x_n}$ dépend de n .

9.2.2 Représentation des poids comme opérations en base n

On note ici $\overline{a_1 \dots a_m}^n = \sum_i a_i n^i$, pour $0 \leq a_i < n$, l'écriture en base n et $\overline{a_1 \dots a_m}^n \rightarrow \overline{b_1 \dots b_p}^n$ la fonction partielle associant à tout entier se décomposant sous la forme $\overline{c_1 \dots c_q a_1 \dots a_m}^n$ l'entier $\overline{c_1 \dots c_q b_1 \dots b_p}^n$. On appelle cette fonction une opération en base n .

Soit $\alpha = ab^*$ le poids d'un chemin dans $\delta\bar{n}$. Écrivons, en oubliant les !, $a = a_1 f_{i_1} a_2 \dots a_k f_{i_k} a_{k+1}$ et $b = b_1 f_{j_1} b_2 \dots b_l f_{j_l} b_{l+1}$ avec aucun des f_i n'apparaissant dans les poids $a_1, b_1, \dots$.

On associe à α l'opération $\overline{j_1 \dots j_l}^n \rightarrow \overline{i_1 \dots i_k}^n$.

A titre d'exemple, au poids de chemin

$$px_3!(f_1x_3!^2(!f_0q)p^*)f_2^*f_0^*)x_3^*p^*$$

de $\delta\bar{3}$ sera associé l'opération $\overline{02^3} \rightarrow \overline{10^3}$. Comme on le verra dans la suite, les opérations associées à des chemins dans $\delta\bar{n}$ seront, au dual près, de la forme :

- $\rightarrow \overline{0\dots 0^n}$, comportant n zéros, dont le chemin associé sera appelé chemin de départ
- $\overline{(n-1)\dots(n-1)^n} \rightarrow$, comportant n fois le chiffre $n-1$, dont le chemin associé sera appelé chemin final (Notons que $\overline{(n-1)\dots(n-1)^n} = n^n - 1$.)
- $\overline{k2\dots 2^n} \rightarrow \overline{(k+1)0\dots 0^n}$ comportant l chiffres 2 et l zéros, pour $0 \leq k < n-1$ et $0 \leq l < n$, dont le chemin associé sera appelé k -ème incrémenteur de rang l .

Par exemple, dans $\delta\bar{3}$, les opérations sont $\rightarrow \overline{000^3}$, $\overline{0^3} \rightarrow \overline{1^3}$, $\overline{1^3} \rightarrow \overline{2^3}$, $\overline{02^3} \rightarrow \overline{10^3}$, $\overline{12^3} \rightarrow \overline{20^3}$, $\overline{022^3} \rightarrow \overline{100^3}$, $\overline{122^3} \rightarrow \overline{200^3}$, $\overline{222^3} \rightarrow$.

9.2.3 Quelques remarques pour faciliter le calcul

Dans δ , ou plutôt dans ce qu'il reste de δ une fois son $\mathfrak{A}$ de tête réduit, on a les deux chemins $\overrightarrow{z_0}$ et $\overleftarrow{z_1}$ correspondant respectivement aux poids z_0q et $z_1p^*z_0^*$.

Un chemin dans $\delta\bar{n}$ commence toujours par $\overrightarrow{z_0}$ et fini toujours par $\overleftarrow{z_0}$. Au milieu, il alterne entre $\overrightarrow{z_1}$ et $\overleftarrow{z_1}$, de sorte que la forme générale d'un chemin soit

$$\overrightarrow{z_0} \square \overrightarrow{z_1} \square \overleftarrow{z_1} \square \dots \square \overrightarrow{z_1} \square \overleftarrow{z_1} \square \overleftarrow{z_0}$$

où les $\square$ doivent être substituées par des chemins dans $\bar{n}$.

On rappelle que f_i et x_i sont de charge i et que $!$ est un morphisme de monoïdes.

9.2.4 Calcul de certains chemins persistants

On pose $\beta(\phi) = x_n!^{n-1}(\phi)f_{n-1}^*$. On a :

Lemme 9.2 $!(\phi)\beta(\psi) = \beta(!(\phi)\psi)$

Preuve

$$!(\phi)\beta(\psi) = !(\phi)x_n!^{n-1}(\psi)f_{n-1}^* = x_n!^{n-1}(!(\phi)\psi)f_{n-1}^*$$

car x_n est de charge n et $!$ est un morphisme. ✎

On utilisera dans la suite la version itérée du lemme précédent :

$$!(\phi)\beta^l(\psi) = \beta^l(!(\phi)\psi)$$

Definition 1 On définit une famille de chemins $(\overrightarrow{\pi_k})$ par :

- $\overrightarrow{\pi_{-1}} = \overrightarrow{z_0} \overleftarrow{x_n} \overrightarrow{z_1}$
- $\forall 0 \leq k \leq n-2$, $\overrightarrow{\pi_k} = \overrightarrow{\pi_{k-1}} \overrightarrow{x_n} \overleftarrow{z_1} \overleftarrow{f_{n-1-k}} \overrightarrow{z_1}$
- $\overrightarrow{\pi_{n-1}} = \overrightarrow{\pi_{n-2}} \overrightarrow{x_n} \overleftarrow{z_1} \overleftarrow{f_0} \overleftarrow{z_0}$

Lemme 9.3 On a les relations suivantes :

$$\begin{aligned} w(\overrightarrow{\pi_{-1}}) &= z_1 f_{n-1}!^{n-1}(p)x_n^*p^* \\ w(\overrightarrow{\pi_k}) &= z_1 f_{n-2-k}!^{n-2-k}(p!(p\beta^{k+1}(p^*)))x_n^*p^* \\ w(\overrightarrow{\pi_{n-1}}) &= qp\beta^n(p^*)x_n^*p^* \end{aligned}$$

Preuve La preuve est un calcul direct. On rappelle que si b est une constante exponentielle de charge c , on peut effectuer les deux réécritures suivantes afin d'aboutir à la forme ab^* :

$$!(x)b \rightarrow b!^c(x) \text{ et } b^*!(x) \rightarrow !^c(x)b^*$$

Commençons par calculer le poids de $\overrightarrow{\pi_{-1}}$.

$$w(\overrightarrow{\pi_{-1}}) = z_1 p^* z_0^* (p f_{n-1} !^{n-1}(p) x_n^* p^* q^*) z_0 q = z_1 f_{n-1} !^{n-1}(p) x_n^* p^*$$

On prouve maintenant les autres relations sur π_k par récurrence sur l . Tout d'abord, on a

$$\begin{aligned} w(\pi_0) &= z_1 p^* z_0^* (p f_{n-2} !^{n-2}(p!(q^*)) f_{n-1}^* p^*) z_0 p z_1^* \\ &\quad ! (q p x_n !^{n-1}(p^*) f_{n-1}^* p^*) z_1 f_{n-1} !^{n-1}(p) x_n^* p^* \\ &= z_1 f_{n-2} !^{n-2}(p!(p x_n !^{n-1}(p^*) f_{n-1}^*)) x_n^* p^* \\ &= z_1 f_{n-2} !^{n-2}(p!(p \beta(p^*))) x_n^* p^* \end{aligned}$$

Si on suppose :

$$w(\pi_k) = z_1 f_{n-2-k} !^{n-2-k}(p!(p \beta^{k+1}(p^*))) x_n^* p^*$$

alors en posant :

$$\tau = w(\overleftarrow{f_{n-2-k} z_1}) = z_1 p^* z_0^* (p f_{n-3-k} !^{n-3-k}(p!(q^*)) f_{n-2-k}^* p^*)$$

on a :

$$\begin{aligned} w(\pi_{k+1}) &= \tau z_0 p z_1^* (q p x_n !^{n-1}(p^*) f_{n-1}^* p^*) z_1 f_{n-2-k} !^{n-2-k}(p!(p \beta^{k+1}(p^*))) x_n^* p^* \\ &= \tau z_0 p f_{n-2-k} !^{n-2-k}(q p x_n !^{n-1}(p^*) f_{n-1}^* ! (p \beta^{k+1}(p^*))) x_n^* p^* \\ &= \tau z_0 p f_{n-2-k} !^{n-2-k}(q p x_n !^{n-1}(\beta^{k+1}(p^*)) f_{n-1}^*) x_n^* p^* \\ &= \tau z_0 p f_{n-2-k} !^{n-2-k}(q p \beta^{k+2}(p^*)) x_n^* p^* \\ &= z_1 f_{n-3-k} !^{n-3-k}(p!(q^*)) !^{n-2-k}(q p \beta^{k+2}(p^*)) x_n^* p^* \\ &= z_1 f_{n-3-k} !^{n-3-k}(p!(p \beta^{k+2}(p^*))) x_n^* p^* \end{aligned}$$

Pour finir, il nous reste à calculer $\overrightarrow{\pi_{n-1}}$:

$$\begin{aligned} w(\overrightarrow{\pi_{n-1}}) &= q^* z_0^* (q q q^* f_0^* p^*) z_0 p z_1^* (q p x_n !^{n-1}(p^*) f_{n-1}^* p^*) z_1 f_0 p ! (p \beta^{n-1}(p^*)) x_n^* p^* \\ &= q p x_n !^{n-1}(p^*) f_{n-1}^* ! (p \beta^{n-1}(p^*)) x_n^* p^* \\ &= q p x_n !^{n-1}(\beta^{n-1}(p^*)) f_{n-1}^* x_n^* p^* \\ &= q p \beta^n(p^*) x_n^* p^* \end{aligned}$$

◀

Soit

$$\nu = w(\overrightarrow{z_1} \overleftarrow{f_0} \overleftarrow{z_1}) = z_0 p z_1^* (p f_0 q q^* q^*) z_1 p^* z_0^* = z_0 p ! (p f_0 q q^* q^*) p^* z_0^*$$

On pose

$$\mu_{k,l} = w\left(\prod_{i=n-l}^{n-l+k} \overrightarrow{f_i} \overrightarrow{z_1} \overleftarrow{f_0} \overleftarrow{z_1}\right) = \prod_{i=n-l}^{n-l+k} \nu w(\overrightarrow{f_i})$$

On a $\mu_{k+1,l} = \nu w(\overrightarrow{f_{n-l+k+1}}) \mu_{k,l}$ ainsi que le lemme suivant.

Lemme 9.4 $\mu_{k,l} = z_0 p f_{n-l+k}!^{n-l} (f_0^k p f_0 q q^*) \tau_l$ où

$$\tau_l = \begin{cases} !^{n-l-1} (p^*) f_{n-l-1}^* p^* z_0^* & \text{si } l < n \\ q^* q^* z_0^* & \text{sinon} \end{cases}$$

Preuve On effectue la preuve par récurrence sur k . Pour $k = 0$, on a, si $l < n$:

$$\begin{aligned} \mu_{0,l} &= \nu w(\overrightarrow{f_{n-l}}) \\ &= z_0 p! (p f_0 q q^* q^*) p^* z_0^*! (p f_{n-l}!^{n-l-1} (! (q) p^*) f_{n-l-1}^* p^*) \\ &= z_0 p! (p f_0 q q^* q^*) f_{n-l}!^{n-l-1} (! (q) p^*) f_{n-l-1}^* p^* z_0^* \\ &= z_0 p f_{n-l}!^{n-l-1} (! (p f_0 q q^*) p^*) f_{n-l-1}^* p^* z_0^* \\ &= z_0 p f_{n-l}!^{n-l} (p f_0 q q^*) \tau_l \end{aligned}$$

et si $l = n$:

$$\begin{aligned} \mu_{0,n} &= \nu w(\overrightarrow{f_0}) \\ &= z_0 p! (p f_0 q q^* q^*) p^* z_0^*! (p f_0 q q^* q^*) \\ &= z_0 p! (p f_0 q q^* q^*) f_0 q q^* q^* z_0^* \\ &= z_0 p f_0 p f_0 q q^* q^* q^* z_0^* \\ &= z_0 p f_0 p f_0 q q^* \tau_n \end{aligned}$$

Supposons maintenant que l'on a obtenu la forme souhaitée pour $\mu_{k,l}$, alors

$$\begin{aligned} \mu_{k+1,l} &= \nu w(\overrightarrow{f_{n-l+k+1}}) \mu_{k,l} \\ &= z_0 p! (p f_0 q q^* q^*) p^* z_0^*! (p f_{n-l+k+1}!^{n-l+k} (! (q) p^*) f_{n-l+k}^* p^*) \\ &\quad z_0 p f_{n-l+k}!^{n-l} (f_0^k p f_0 q q^*) \tau_l \\ &= z_0 p! (p f_0 q q^* q^*) f_{n-l+k+1}!^{n-l+k} (! (q) p^*) !^{n-l} (f_0^k p f_0 q q^*) \tau_l \\ &= z_0 p f_{n-l+k+1}!^{n-l+k} (! (p f_0 q q^*) p^*) !^{n-l} (f_0^k p f_0 q q^*) \tau_l \\ &= z_0 p f_{n-l+k+1}!^{n-l} (!^k (! (p f_0 q q^*) p^*) f_0^k p f_0 q q^*) \tau_l \\ &= z_0 p f_{n-l+k+1}!^{n-l} (f_0^k ! (p f_0 q q^*) f_0 q q^*) \tau_l \\ &= z_0 p f_{n-l+k+1}!^{n-l} (f_0^{k+1} p f_0 q q^*) \tau_l \end{aligned}$$

✎

Théorème 9.5 Il y a **au moins** $4 + 2n(n-1)$ chemins dans $\delta\bar{n}$. En d'autres termes, $|\delta\bar{n}| \geq 2 + n(n-1)$.

Nommé on a

- le chemin de départ

$$\overrightarrow{S_n} = \overrightarrow{z_0} (\prod_{i=0}^{n-1} \overrightarrow{f_i z_1} \overrightarrow{f_0 z_1}) \overleftarrow{z_0}$$

de poids $\sigma_n = p x_n f_0^n q q^* q^*$

- le k -ème incrémenteur de rang l , pour $k < n-1$, $l \leq n-1$,

$$\overrightarrow{I_n(k,l)} = \overrightarrow{\pi_{l-1} f_{k+1} z_1} (\prod_{i=n-l}^{n-1} \overrightarrow{f_i z_1} \overrightarrow{f_0 z_1}) \overrightarrow{x_n z_0}$$

de poids

$$\iota_n(k, l) = px_n!^{n-l-1}(f_{k+1}!^k(\beta^l(!f_0^l q)p^*))f_k^*x_n^*p^*$$

– le chemin final $\overrightarrow{E_n} = \overrightarrow{\pi_{n-1}}$ de poids $\epsilon_n = qp\beta^n(p^*)x_n^*p^*$
ainsi que leurs inversions d'orientation.

Preuve

– Avec les définitions précédentes de $\mu_{k,l}$, on a

$$\begin{aligned} w(\overrightarrow{S_n}) = \sigma_n &= w(\overrightarrow{x_n z_0})\mu_{n-1,n}w(\overrightarrow{z_0}) \\ &= q^*z_0^!(qp x_n!^{n-1}(p^*)f_{n-1}^*p^*)z_0 p f_{n-1} f_0^{n-1} p f_0 q q^* q^* q^* z_0^* z_0 q \\ &= p x_n!^{n-1}(p^*)f_0^{n-1} p f_0 q q^* q^* \\ &= p x_n f_0^n q q^* q^* \end{aligned}$$

– Pour I_n , on a :

$$\begin{aligned} w(\overrightarrow{I_n(k, l)}) = \iota_n(k, l) &= w(\overrightarrow{x_n z_0})\mu_{l-1,l}w(\overrightarrow{\pi_{l-1} f_{k+1} z_1}) \\ &= w(\overrightarrow{x_n z_0})\mu_{l-1,l}z_0 p z_1^!(p f_{k+1}!^k(!q)p^*)f_k^*p^* \\ &\quad z_1 f_{n-l-1}!^{n-l-1}(p!(p\beta^l(p^*)))x_n^*p^* \\ &= w(\overrightarrow{x_n z_0})\mu_{l-1,l}z_0 p f_{n-l-1}!^{n-l-1}(p f_{k+1}!^k(!q)p^*)f_k^*(p\beta^l(p^*)))x_n^*p^* \\ &= w(\overrightarrow{x_n z_0})\mu_{l-1,l}z_0 p f_{n-l-1}!^{n-l-1}(p f_{k+1}!^k(!q)\beta^l(p^*))f_k^*x_n^*p^* \\ &= w(\overrightarrow{x_n z_0})z_0 p f_{n-1}!^{n-l-1}(!f_0^{l-1} p f_0 q q^*)f_{k+1}!^k(!q)\beta^l(p^*))f_k^*x_n^*p^* \\ &= w(\overrightarrow{x_n z_0})z_0 p f_{n-1}!^{n-l-1}(f_{k+1}!^k(!f_0^{l-1} p f_0 q)\beta^l(p^*))f_k^*x_n^*p^* \\ &= w(\overrightarrow{x_n z_0})z_0 p f_{n-1}!^{n-l-1}(f_{k+1}!^k(\beta^l(!f_0^{l-1} p f_0 q)p^*))f_k^*x_n^*p^* \\ &= p x_n!^{n-1}(p^*)!^{n-l-1}(f_{k+1}!^k(\beta^l(!f_0^{l-1} p f_0 q)p^*))f_k^*x_n^*p^* \\ &= p x_n!^{n-l-1}(f_{k+1}!^k(!p^*)\beta^l(!f_0^{l-1} p f_0 q)p^*)f_k^*x_n^*p^* \\ &= p x_n!^{n-l-1}(f_{k+1}!^k(\beta^l(!p^*)f_0^{l-1} p f_0 q)p^*)f_k^*x_n^*p^* \\ &= p x_n!^{n-l-1}(f_{k+1}!^k(\beta^l(!f_0^l q)p^*)f_k^*)x_n^*p^* \end{aligned}$$

– Dans la mesure où $\overrightarrow{E_n} = \overrightarrow{\pi_{n-1}}$, on a déjà le poids qui lui correspond.

◀

On peut vérifier sur les poids que les opérations associées sont bien celles définies au paragraphe 9.2.2.

9.2.5 Complétude d'un ensemble de chemins

Maintenant qu'on a de bons candidats pour la GdI, nous aimerions nous assurer que ce sont les seuls chemins.

Rappelons tout d'abord une propriété d'adéquation due à Danos et Regnier :

Lemme 9.6 Soit R un réseau de SMELL, $\forall \phi, \phi' \in \mathfrak{Pr}_f(R)$, $\phi \neq \phi' \iff w(\phi)w(\phi')^* = 0$

Si $\{\phi_1, \dots, \phi_k\} \subseteq \mathfrak{Pr}_f(R)$, chaque ϕ_i a un poids qui peut être mis sous forme ab^* , on pose donc $w(\phi_i) = a_i b_i^*$, et on a donc

$$\begin{aligned} \left(\sum_i w(\phi_i) + w(\phi_i)^*\right) \left(\sum_i w(\phi_i) + w(\phi_i)^*\right)^* &= \sum_i w(\phi_i) w(\phi_i)^* + \sum_i w(\phi_i)^* w(\phi_i) \\ &= \sum_i a_i a_i^* + b_i b_i^* \end{aligned}$$

On en déduit le lemme suivant qui nous permettra d'affirmer qu'un ensemble de chemins est complet.

Lemme 9.7 *Soient R un réseau de SMELL, M un modèle non-trivial de M_R tel que et $X \subseteq \mathfrak{Pr}_f(R)$ tel que*

- X est clos par inversion d'orientation
- et $M \models \text{Ex}_X \text{Ex}_X^* = 1$ où $\text{Ex}_X = \sum_{\phi \in X} w(\phi)$.

Alors on a $X = \mathfrak{Pr}_f(R)$.

9.2.6 Branchements d'arbres complets

Notons X_n l'ensemble des chemins données par le théorème 9.5 appliqué à $\delta\bar{n}$. Dans cette partie on va examiner une condition nécessaire pour que X_n satisfasse les hypothèses du lemme précédent.

9.2.6.a Le cas $n = 3$

En appliquant le théorème 9.5 on a, pour $n = 3$, au moins huit chemins. On donne dans le tableau suivant leur décomposition en forme ab^* présenté en forme canonique à droite ¹

	a	b
S_3	$p!(q)x_3!^2(f_0)!(f_0)f_0$	qq
$I_3(0,0)$	$p!(q)x_3!^2(f_1)$	$p!(p)x_3!^2(f_0)$
$I_3(1,0)$	$p!(q)x_3!^2(f_2)$	$p!(p)x_3!^2(f_1)$
$I_3(0,1)$	$p!(q)x_3!^2(f_0)!(f_1)!(x_3)$	$p!(p)x_3!^2(f_2)!(f_0)$
$I_3(1,1)$	$p!(q)x_3!^2(f_0)!(f_2)!^2(x_3)$	$p!(p)!x_3!^2(f_2)!(f_1)$
$I_3(0,2)$	$p!(q)x_3!^2(f_0)!(f_0)f_1x_3!^2(x_3)$	$p!(p)x_3!^2(f_2)!(f_2)f_0$
$I_3(1,2)$	$p!(q)x_3!^2(f_0)!(f_0)f_2!(x_3!^2(x_3))$	$p!(p)x_3!(f_2)f_2f_1$
E_3	$qp x_3!^2(x_3!^2(x_3))$	$p!(p)x_3!(f_2)f_2f_2$

La figure 9.1 donne une superposition de ces poids en présentant l'arbre des préfixes. Pour calculer la valeur de $\text{Ex}_{X_3} \text{Ex}_{X_3}^*$ il suffit de *brancher* un arbre dual en face où tous les poids ont été étoilés, puis de faire la somme de toutes les chemins de la racine du premier à la racine du second. Notons par exemple que si l'on sait que

$$qq^* + p x_3!^2(x_3!^2(x_3 x_3^*) x_3^*) x_3^* p^* = 1$$

on peut simplifier le calcul en récupérant un nouveau branchement où l'on a remplacé

¹c'est-à-dire la forme correspondant à pousser au maximum les constantes exponentielles sur la droite.

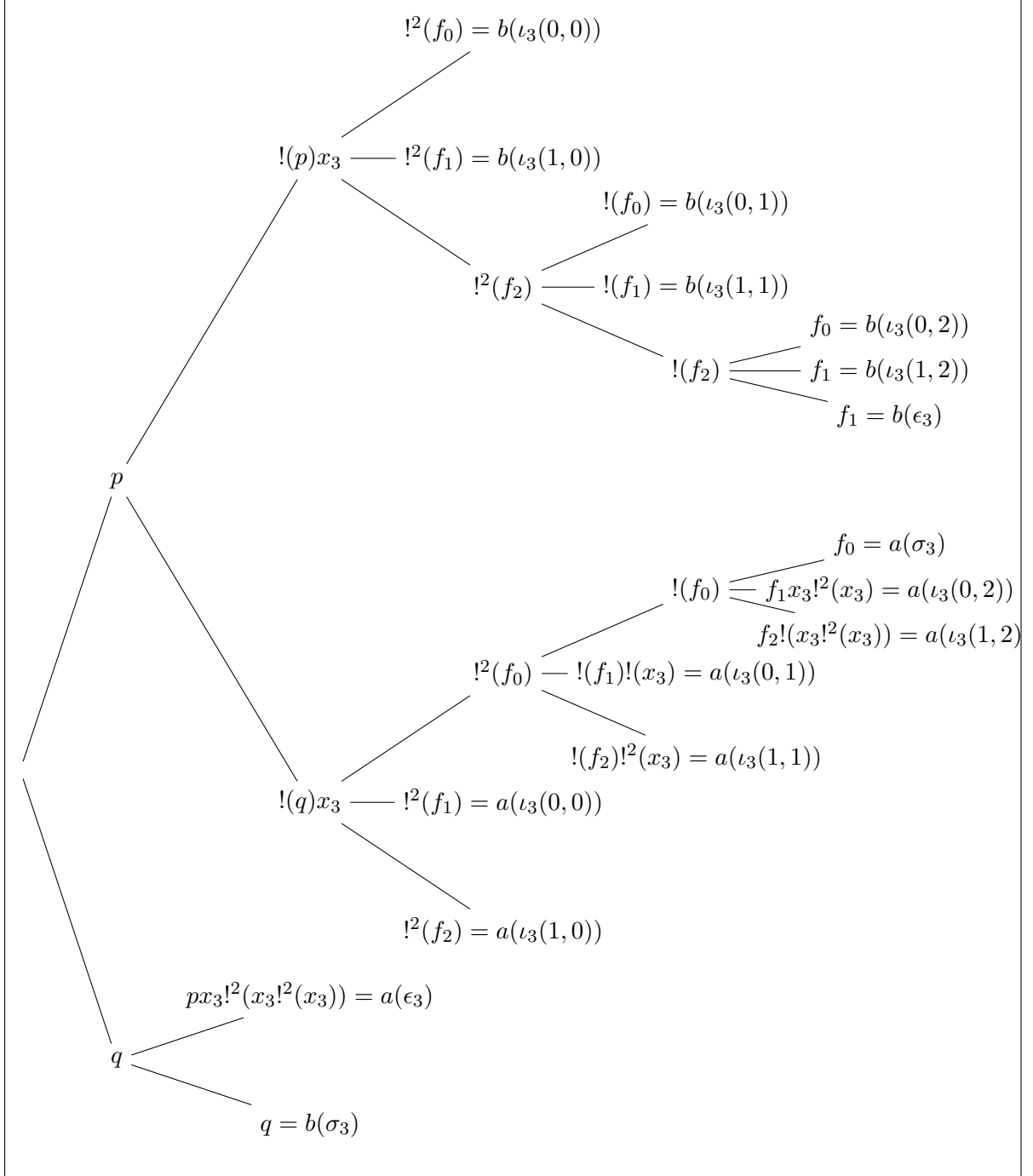
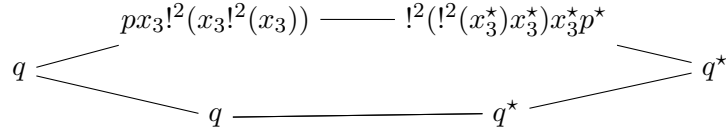


Figure 9.1: Superposition des poids de chemins dans X_3



par $q \longrightarrow q^*$.

9.2.6.b Réduction d'arbres complets

On va maintenant essayer de généraliser le calcul précédent. Tout d'abord supposons que l'on connaisse un modèle M ayant la propriété de partitionnement pour un réseau R , c'est-à-dire tel que pour tout paquet de constantes exponentielles $(\beta_i)_i$ apparaissant dans R on ai

$$!^h(\sum_i \beta_i \beta_i^*) = 1$$

Lemme 9.8 *Soit $(\pi_n)_n$ un ensemble de poids de chemins de $\mathfrak{Pr}_f(R)$ tel que pour tout paquet de constantes exponentielles $(\beta_i)_i$ (ou pour $\{p, q\}$) si $\exists \alpha, i_0$ avec $\alpha !^h(\beta_{i_0})$ préfixe d'au moins un des π_n , alors $\alpha !^h(\beta_i)$ est également un préfixe pour tout i .*

On a $M \models \sum_n \pi_n \pi_n^* = 1$.

La propriété des π_n est appelé *avoir une superposition complète*. On vérifie directement que X_3 a une superposition complète.

9.2.6.c Un modèle ayant la propriété de partitionnement

On montre que le modèle des injections partielles de $\mathbb{N}$ a la propriété de partitionnement pour $\delta\bar{n}$. On rappelle que p et q sont modélisés par $p(n) = 2n$ et $q(n) = 2n + 1$. Fixons une bijection $\lceil \cdot, \cdot \rceil : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{N}$, on définit $(u \otimes v)(\lceil x, y \rceil) = \lceil u(x), v(y) \rceil$. Posons $!(u) = 1 \otimes u$. Si $\tau : \mathbb{N}^n \rightarrow \mathbb{N}$ est une injection partielle, on la projette sur les injections partielles de $\mathbb{N}$ par

$$u_\tau(\lceil x_1, \dots, x_{n+1} \rceil) = \lceil \tau(x_1, \dots, x_n), x_{n+1} \rceil$$

Alors, si τ est de domaine plein, on a $u_\tau^* u_\tau = 1$, si τ est à la fois de domaine plein et d'image pleine on a $u_\tau u_\tau^* = 1$ et si τ et τ' sont disjointes $u_\tau^* u_{\tau'} = 0$. De plus u_τ satisfait l'équation

$$!(v)u_\tau = u_\tau !^n(v)$$

Il nous faut alors trouver des τ appropriés pour chaque constante exponentielle dans $\delta\bar{n}$. Pour modéliser x_n de charge n on pose $x_n = u_{\xi_n}$ avec $\xi_n(x_1, \dots, x_n) = \lceil x_1, \dots, x_n \rceil$, cela nous donne $x_n = 1$. Ensuite, on définit $\phi_k : \mathbb{N}^k \rightarrow \mathbb{N}$ par $\phi_0 = 0$, $\phi_{n-1}(x_1, \dots, x_{n-1}) = 2^{n-1} \lceil x_1, \dots, x_{n-1} \rceil$ et

$$\phi_k(x_1, \dots, x_k) = 2^{k-1}(2 \lceil x_1, \dots, x_k \rceil + 1)$$

Et on pose $f_k = u_{\phi_k}$, il est alors clair que $\sum_k f_k f_k^* = 1$.

Ainsi, on a construit un modèle approprié avec propriété de partitionnement pour $\delta\bar{n}$.

9.2.6.d Une superposition complète de X_n

Calculons les formes canoniques à droite des poids donnés par le théorème 9.5. On va juste détailler le calcul pour $a(\sigma_n)$. On a

$$\begin{aligned} a(\sigma_n) &= px_n f_0^n q \\ &= px_n !^n(q) f_0^n \\ &= p!(q) x_n f_0^{n-1} f_0 \\ &= p!(q) x_n !^{n-1}(f_0) f_0^{n-1} \\ &= p!(q) x_n !^{n-1}(f_0) !^{n-2}(f_0) \dots f_0 \end{aligned}$$

Plus généralement on trouve :

$$\begin{aligned} a(\sigma_n) &= p!(q) x_n \prod_{i=0}^{n-1} !^{n-1-i}(f_0) \quad b(\sigma_n) = qq \\ a(\epsilon_n) &= qp \prod_{i=0}^{n-1} !^{i(n-1)}(x_n) \quad b(\epsilon_n) = p!(p) x_n \prod_{i=0}^{n-1} !^{n-1-i}(f_{n-1}) \\ a(\iota_n(k, l)) &= px_n !^{n-l-1}(f_{k+1} !^k(\prod_{i=0}^{l-1} !^{i(n-1)}(x_n) !^{l(n-1)+1}(f_0^l q))) \\ &= px_n !^{n-l-1}(f_{k+1} !^{k+1}(f_0^l q) !^k(\prod_{i=0}^{l-1} !^{i(n-1)}(x_n))) \\ &= px_n !^{n-l-1}(!^k(f_0^l q) f_{k+1} !^k(\prod_{i=0}^{l-1} !^{i(n-1)}(x_n))) \\ &= px_n !^n(q) (\prod_{i=0}^{l-1} !^{n-i-1}(f_0)) !^{n-l-1}(f_{k+1} !^k(\prod_{i=0}^{l-1} !^{i(n-1)}(x_n))) \\ &= p!(q) x_n (\prod_{i=0}^{l-1} !^{n-i-1}(f_0)) !^{n-l-1}(f_{k+1} !^k(\prod_{i=0}^{l-1} !^{i(n-1)}(x_n))) \\ b(\iota_n(k, l)) &= px_n !^{n-l-1}(f_k !^k(\prod_{i=0}^{l-1} !^{i(n-1)}(f_{n-1}) !^{l(n-1)}(p))) \\ &= px_n !^{n-l-1}(!^k(\prod_{i=0}^{l-1} !^{i(n-1)}(f_{n-1}) !^{l(n-1)}(p)) f_k) \\ &= px_n !^{n-l-1}(!^{l+1}(p) \prod_{i=0}^{l-1} !^{i(n-1)+1}(f_{n-1}) f_k) \\ &= p!(p) x_n !^{n-l-1}(\prod_{i=0}^{l-1} !^{i(n-1)+1}(f_{n-1}) f_k) \\ &= p!(p) x_n \prod_{i=0}^{l-1} !^{n-i-1}(f_{n-1}) !^{n-l-1}(f_k) \end{aligned}$$

On a donc une superposition très proche de celle du cas $n = 3$, on montre directement qu'elle est complète.

9.2.6.e Le théorème final

En mettant bout à bout ce que nous avons montré aux paragraphes précédents, on obtient le théorème suivant :

Théorème 9.9 $|\delta\bar{n}| = 2 + n(n - 1)$

Rappelons que $|\bar{n^n}| = 1 + n^n$, on a donc un gain spectaculaire.

9.3 Conjectures

Remarquons que le calcul précédent repose presque entièrement sur la définition des opérations en base n associées aux chemins. Nous donnons ici, sans les prouver, les opérations associées aux chemins correspondant à d'autres calculs sur les entiers de Church.

Pour $t = \bar{n}$ on a $|t| = n + 1$ et les opérations sont

$$\begin{array}{rcl} & \rightarrow & \bar{0}^n \\ \bar{0}^n & \rightarrow & \bar{1}^n \\ & \vdots & \\ \overline{n-2}^n & \rightarrow & \overline{n-1}^n \\ \overline{n-1}^n & \rightarrow & \end{array}$$

Pour $t = \bar{+} \bar{n} \bar{m}$ on a $|t| = 1 + n + m$ et les opérations sont

$$\begin{array}{rcl} & \rightarrow & \bar{0}^n \\ \bar{0}^n & \rightarrow & \bar{1}^n \\ & \vdots & \\ \overline{n-2}^n & \rightarrow & \overline{n-1}^n \\ \overline{n-1}^n & \rightarrow & \bar{0}^m \\ \bar{0}^m & \rightarrow & \bar{1}^m \\ & \vdots & \\ \overline{m-2}^m & \rightarrow & \overline{m-1}^m \\ \overline{m-1}^m & \rightarrow & \end{array}$$

Pour $t = \overline{\times} \overline{n} \overline{m}$ on a $|t| = 1 + n + m$ et les opérations sont

$$\begin{array}{rcl}
 & & \rightarrow \overline{0^m 0^n} \\
 \overline{0^n} & \rightarrow & \overline{1^n} \\
 & \vdots & \\
 \overline{n - 2^n} & \rightarrow & \overline{n - 1^n} \\
 \overline{0^m n - 1^n} & \rightarrow & \overline{1^m 0^n} \\
 \overline{1^m n - 1^n} & \rightarrow & \overline{2^m 0^n} \\
 & \vdots & \\
 \overline{m - 2^m n - 1^n} & \rightarrow & \overline{m - 1^m 0^n} \\
 \overline{m - 1^m n - 1^n} & \rightarrow &
 \end{array}$$

où $\overline{a^m b^n} = b + na$.

Pour $t = \overline{m} \overline{n}$ on a $|t| = 2 + m(n - 1)$ et les opérations sont

$$\begin{array}{rcl}
 & \xrightarrow{m} & \overline{0 \dots 0^n} \\
 \overline{0^n} & \xrightarrow[1]{1} & \overline{1^n} \\
 & \vdots & \\
 \overline{n - 2^n} & \xrightarrow[1]{2} & \overline{0(n - 1)^n} \\
 \overline{0(n - 1)^n} & \xrightarrow[2]{2} & \overline{1(n - 1)^n} \\
 & \vdots & \\
 \overline{(n - 2)(n - 1)^n} & \xrightarrow[2]{3} & \overline{0(n - 1)(n - 1)^n} \\
 & \vdots & \\
 \overline{0(n - 1) \dots (n - 1)^n} & \xrightarrow[m]{m} & \overline{1(n - 1) \dots (n - 1)^n} \\
 & \vdots & \\
 \overline{(n - 2)(n - 1) \dots (n - 1)^n} & \xrightarrow[m]{m} & \overline{(n - 1)(n - 1) \dots (n - 1)^n} \\
 \overline{(n - 1) \dots (n - 1)^n} & \xrightarrow{m} &
 \end{array}$$

où l'on a indiqué par $\xrightarrow[a]{b}$ que l'opération prenait un nombre à a chiffres et renvoyé un nombre à b chiffres.

Nous avons par exemple $|\overline{2025}| = 2026$ mais $2025 = 3^4 5^2$ et on obtient un $|t| = 1 + (2 + 4 * (3 - 1)) + (2 + 2 * (5 - 1)) = 21$ avec $t \rightarrow_{\beta} \overline{2025}$.

Troisième partie

Les Réseaux d'Interaction Différentiels et leur Géométrie de l'Interaction

Chapitre 10

Réseaux d'interaction différentiels

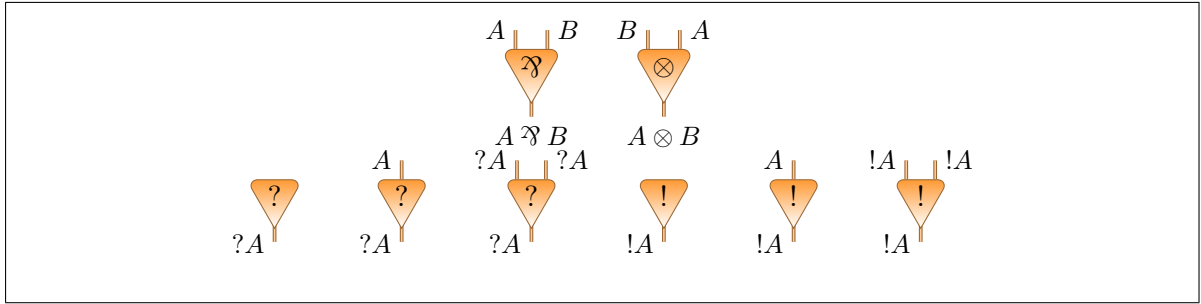


Figure 10.1: Les cellules des réseaux d'interaction différentiels et leur typage

Ce chapitre reprend essentiellement les définitions et résultats de l'article fondateur de Ehrhard et Regnier [ER05].

10.1 Définitions

10.1.1 Réseaux simples

Les cellules des réseaux d'interactions et leur typage sont présentés sur la figure 10.1. Les cellules sont donc une symétrisation de MELL sans la promotion. Ces nouvelles cellules sont naturellement appelées *co-déréliciton*, *co-contraction* et *co-affaiblissement*.

Un réseau d'interaction ainsi construit sera appelé *réseau d'interaction différentiel simple*, ou juste réseau simple.

On note $\mathfrak{R}_\partial^0$ l'ensemble des réseaux simples.

10.1.2 Réseaux d'interaction différentiels

Soit $P \subseteq \mathbb{N}$, on note $\mathfrak{R}\langle P \rangle = \{R \in \mathfrak{R}_\partial^0 \mid P_f(R) = P\}$.

On pose $\mathfrak{R}_\partial = \bigcup_{P \subseteq \mathbb{N}} \mathbb{N}[\mathfrak{R}\langle P \rangle]$, qui est une union disjointe de $\mathbb{N}$ -modules libres.

Un élément de $\mathfrak{R}_\partial$ est appelé réseau d'interaction différentiel, il est par définition de la forme $\sum_i \lambda_i R_i$ où les $\lambda_i \in \mathbb{N}$ et les R_i sont des réseaux simples ayant tous les mêmes ports libres P , par extension on dira que le réseau lui-même a pour port libres P .

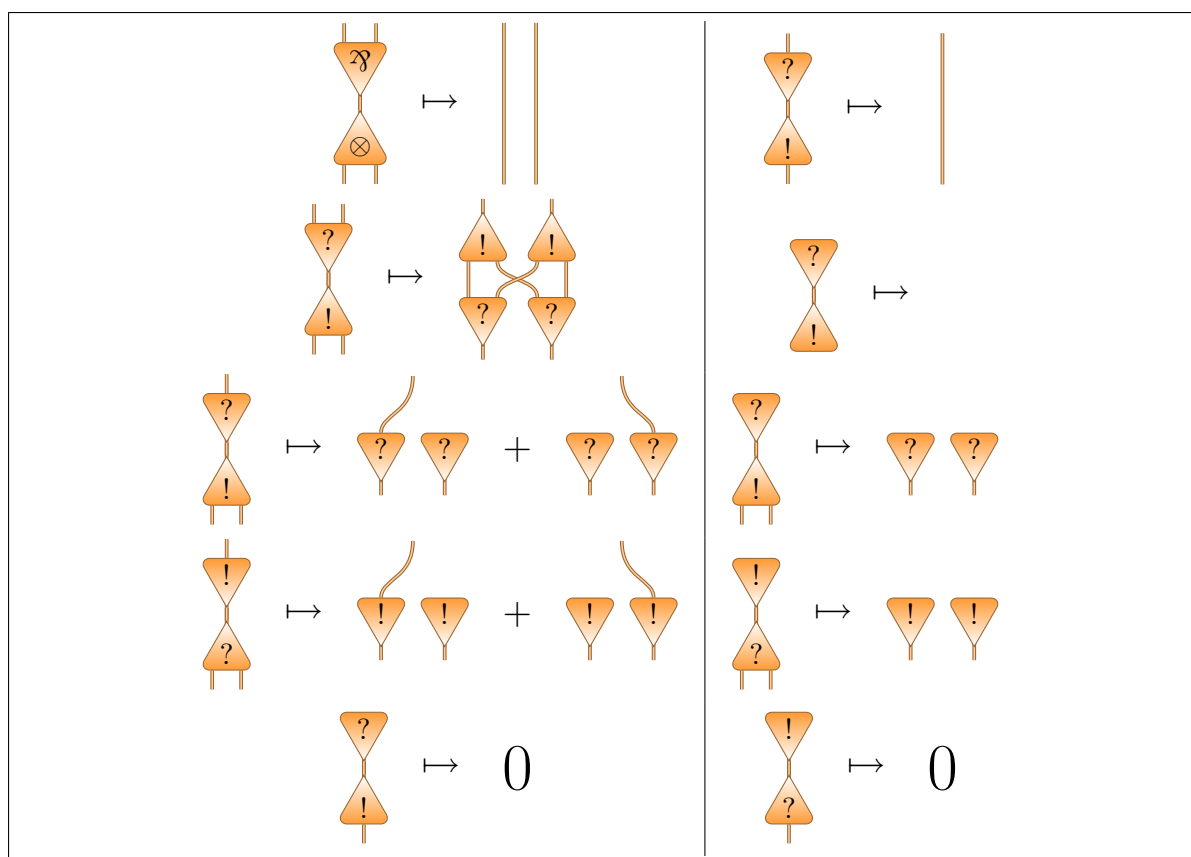


Figure 10.2: Les règles de reduction des réseaux d'interaction différentiels

Remarque 10.1 *La présence des coefficients est superflue. En effet, on peut toujours remplacer λR par la somme $R + \dots + R$ comportant λ termes. Cette remarque n'a pas grand intérêt du point de vue des modules, mais elle nous permettra, d'une part d'alléger les notations, et d'autre part de définir la réduction de manière atomique.*

Pour chaque ensemble $P \subseteq \mathbb{N}$ on note 0_P le zéro du module $\mathbb{N}[\mathfrak{R}\langle P \rangle]$.

Par abus de notation on utilisera souvent 0 en rendant implicite le type, de même on aura tendance à considérer les réseaux d'interaction différentiels comme le $\mathbb{N}$ -module libre engendré par les réseaux simples. Cela constitue a priori une approximation grossière, cependant, cela sera vérifié du point de vue de la sémantique.

Dans la suite, la notation $R + R'$ signifiera implicitement que R et R' ont les mêmes ports libres.

10.1.3 Interfaces

Soient $R + \sum_i R_i$ un réseau d'interaction différentiels et I une interface contenu dans R , alors chacun des R_i contient également I . Et il semble naturel de considérer que la somme elle-même contient cette interface.

Définition 10.2 *Soit $R = \sum_i R_i$ un réseau d'interaction différentiel, on dit que R contient l'interface $I \subseteq \mathbb{N}$ si un des R_i la contient.*

On peut donc étendre la composition des réseaux d'interactions aux sommes, on pose $\sum_i R_i^I \rightsquigarrow R'^I = \sum_i (R_i^I \rightsquigarrow R'^I)$.

10.1.4 Réduction dans un $\mathbb{N}$ -module libre

Soit R un ensemble et $\rightarrow$ une relation de $R \times \mathbb{N}[R]$. On étend $\rightarrow$ à $\mathbb{N}[R] \times \mathbb{N}[R]$ par

$$\sum_i r_i \rightarrow_+ \sum_{i \neq i_0} r_i + \sum_k r'_k \iff \exists i_0, r_{i_0} \rightarrow \sum_k r'_k$$

Par abus de notation on utilisera juste $\rightarrow$ pour faire référence à $\rightarrow_+$ lorsqu'il n'y aura pas d'ambiguïté.

Définition 10.3 *Une relation $\rightarrow$ de $R \times \mathbb{N}[R]$ est dite n -large si et seulement si*

$$\forall r \in R, \sum_{i=1}^m r_i \in \mathbb{N}[R], r \rightarrow \sum_{i=1}^m r_i \Rightarrow m \leq n$$

La relation est dite fortement n -large si l'inégalité $m \leq n$ est une égalité pour au moins un couple.

La relation est dite fortement confluente terme à terme si et seulement si

- pour tout $r \in R$ tel que $r \rightarrow \sum_{i=1}^n r_i$ et $r \rightarrow \sum_{j=1}^m r'_j$, on a $n = m$ et pour tout i, j , il existe r' tel que $r_i \rightarrow r'$ et $r'_j \rightarrow r'$,
- pour tout $r \in R$ tel que $r \rightarrow \sum_i r_i$ et $r \rightarrow 0$, et pour tout i on a $r_i \rightarrow 0$

Fait 10.4 *Si $\rightarrow$ est n -large et fortement confluente terme à terme, alors $\rightarrow_+$ est n -confluente, c'est-à-dire que pour tout $r \in \mathbb{N}[R]$ tel que $r \rightarrow_+ r_1$ et $r \rightarrow_+ r_2$, avec $r_1 \neq r_2$, il existe $r' \in \mathbb{N}[R]$ tel que $\forall i \in \{1, 2\}, r_i \rightarrow_+^m r'$ pour un certain $m \leq n$.*

Preuve Supposons que $r_1 = \sum_{i=1}^m r_{1,i}$ et que $r_2 = 0$. Par hypothèse, on a $\forall i, r_{1,i} \rightarrow 0$ et donc

$$\sum_{i=1}^m r_{1,i} \rightarrow_+ 0 + \sum_{i=2}^m r_{1,i} \rightarrow_+ \cdots \rightarrow_+ 0$$

soit, de manière concise $\sum_i r_{1,i} \xrightarrow{+}_m 0$. Or, $0 \xrightarrow{+}_0 0$.

Supposons maintenant que $r_1 = \sum_{i=1}^m r_{1,i}$ et $r_2 = \sum_{i=1}^m r_{2,i}$. Par hypothèse, pour tout i , il existe r'_i tel que $r_{1,i} \rightarrow r'_i$ et $r_{2,i} \rightarrow r'_i$. On a alors, pour $k \in \{1, 2\}$:

$$\sum_{i=1}^m r_{k,i} \rightarrow_+ r'_1 + \sum_{i=2}^m r_{k,i} \rightarrow_+ \cdots \rightarrow_+ \sum_i r'_i$$

Et on conclut en posant $r' = \sum_i r'_i$. ◀

10.1.5 Réduction des réseaux d'interaction différentiels

On définit une famille de réductions $(\rightarrow_P)_{P \subseteq \mathbb{N}}$, telle que $\rightarrow_P \subseteq \mathfrak{R}\langle P \rangle \times \mathbb{N}[\mathfrak{R}\langle P \rangle]$ soit engendrée par les règles présentées dans la figure 10.2.

Deux règles (et leurs duales), celles faisant apparaître à droite la structure de $\mathbb{N}$ -module, nécessitent d'être explicitées :

si R est un réseau contenant un redex $!^1/?^2$, alors R est de la forme $R' \perp_f r$ où r est le redex et f un raccord de l'interface de r . La règle nous donne $r \rightarrow r_1 + r_2$, et on définit $R \rightarrow R' \perp_f r_1 + R' \perp_f r_2$. Ce réduit est bien défini car l'interface est préservée par réduction.

si R est un réseau contenant un redex $!^0/?^1$, on a également R de la forme $R' \perp_f r$ où r est le redex. La règle nous donne $r \rightarrow 0_{T_r}$ où T_r est le type de r , et on définit $R \rightarrow 0_T$ où T est le type de R .

Ces deux règles sont donc fortement globales : elles ne procèdent pas d'une réécriture locale de réseau et de son passage au contexte. Dans le premier cas, il y a effectivement duplication de part et d'autre de la somme du réseau autour du redex. Il semble cependant que cela n'est que la compatibilité avec la somme d'une réduction locale. En revanche, dans le second cas le réseau autour est complètement détruit par la réduction.

La réduction ainsi défini est fortement 2-large et fortement confluente terme à terme, puisque vis-à-vis des termes tout se passe dans les réseaux d'interaction. Ainsi on a le fait suivant :

Fait 10.5 *la réduction des réseaux d'interaction différentiels est 2-confluente.*

10.1.6 Réseaux différentiels bien typés

Définition 10.6 *On dit que $R = \sum_i \lambda_i R_i \in \mathfrak{R}\langle P \rangle$ est bien typé si et seulement si*

- *chacun des R_i est bien typé,*
- *et que le type de $p \in P$ est le même dans tous les R_i .*

On parlera alors d'interface typée et de types pour R .

10.2 Interprétations des règles de réduction

Nous allons donner ici différentes interprétations des cellules et règles de réduction des réseaux d'interaction différentiels. Ce sont des intuitions et non pas des définitions réelles qui vont être données, cependant, elles peuvent se transcrire à chaque fois en des définitions rigoureuses. En effet, pour la première interprétation, il existe effectivement une sémantique dans laquelle la co-déréliction correspond à la dérivation, et pour la seconde interprétation il existe, on le verra au chapitre suivant, un codage d'un calcul de processus dans les réseaux d'interaction différentiels.

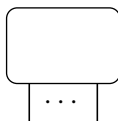
10.2.1 Dérivation de fonctions analytiques

Cette interprétation correspond à l'approche originelle de Ehrhard et Regnier que l'on peut trouver dans l'article [ER05].

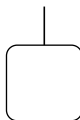
On se place désormais dans un $\mathbb{C}^n$.

10.2.1.a Interprétation des réseaux

Nous allons considérer ici une orientation verticale des réseaux. Un réseau de la forme



représente une fonction $\mathbb{C}^n \mapsto \mathbb{C}$, envoyant $(0, \dots, 0)$ sur 0, alors qu'un réseau de la forme



représente un argument que l'on pourra brancher sur une fonction.

Du point de vue du typage nous allons considérer un seul niveau de modalité exponentielle. Les types de base correspondent à des domaines dans $\mathbb{C}^n$. Une variable typée par une formule A sans exponentielle signifiera que la fonction est linéaire en cette variable, réciproquement une variable typée $?A$ signifiera que la fonction est analytique en cette variable. Par dualité, un argument de type A devra être utilisé exactement une fois, alors qu'un argument de type $!A$ sera utilisable un nombre quelconque de fois.

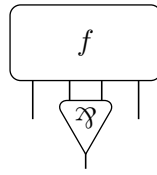
10.2.1.b Interprétation de la somme

La somme est ici naturellement interprétée comme la somme usuelle de l'espace vectoriel considéré.

10.2.1.c Interprétation des cellules

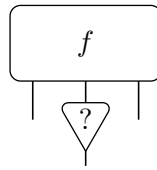
Les cellules vont maintenant être séparées en deux groupes : les cellules permettant de construire des fonctions et les cellules permettant de construire des arguments.

Passage d'arguments en couple Si $f(x, y, \dots)$ est une fonction à n arguments, le réseau



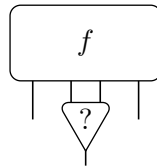
représente la fonction à $n - 1$ arguments $g((x, y), \dots) = f(x, y, \dots)$. Ces couples ne vivent plus dans $\mathbb{C}$ et cela sort de notre cadre, mais pour l'utilisation que l'on en fera, cela n'est pas si embêtant.

Fonction linéaire vue comme analytique Si $f(x, \dots)$ est une fonction à n arguments linéaire en x , le réseau



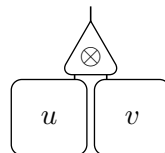
représente la même fonction considérée comme analytique en x .

Passage d'arguments similaires Si $f(x, y, \dots)$ est une fonction à n arguments analytique en x et y , le réseau



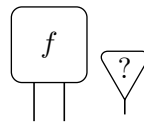
représente la fonction à $n - 1$ arguments $g(x, \dots) = f(x, x, \dots)$ également analytique en x . Notons que justement cette construction ne préserve pas la linéarité.

Couples d'arguments Le réseau



représente l'argument couple (u, v) .

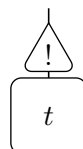
Fonction constante Si f est une fonction à n arguments, le réseau



représente la fonction à $n + 1$ arguments $g(x, y_1, \dots, y_n) = f(y_1, \dots, y_n)$, c'est-à-dire une fonction analytique constante vis-à-vis de l'argument x .

Lorsqu'il n'y a pas d'autres arguments on obtient la fonction nulle.

Dérivation en 0 Un argument de la forme



correspond à un argument t utilisable une seule fois devant être passé à une fonction analytique. Si on branche cette argument sur la variable x d'une fonction $f(x, y, \dots)$ cela correspond donc à évaluer une approximation linéaire de f sur la variable x en t . En effectuant un développement limité à l'ordre 1, on a :

$$f(x_0 + t, y, \dots) = f(x_0, y, \dots) + t \times \frac{\partial f}{\partial x}(x_0, y, \dots) + o(t)$$

que l'on peut simplifier avec $x_0 = 0$ dans notre cas en

$$f(t, y, \dots) = t \times \frac{\partial f}{\partial x}(0, y, \dots) + o(t)$$

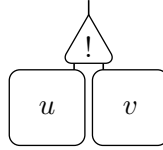
La co-déréliction correspond alors à la fonction $(t, y, \dots) \mapsto t \times \frac{\partial f}{\partial x}(0, y, \dots)$.

Argument zéro Le réseau



correspond au plus simple argument utilisable en un nombre quelconque de fois : 0.

Superposition d'arguments Le réseau



représente la superposition des arguments u et v , c'est-à-dire l'argument $u + v$.

10.2.1.d Interprétation des règles

$\mathcal{A}/\otimes$ Il s'agit de calculer $g((u, v), \dots) = f(u, v, \dots)$ et donc le résultat est juste la fonction f .

$?_0/!_0$ On applique une fonction constante en 0, tout se passe comme si on réduisait le nombre de variables.

$?_0/!_1$ On dérive une fonction à n variables sur une variable pour laquelle la fonction est constante, cela donne la fonction nulle à $n - 1$ variables.

$?_0/!_2$ On calcule $f(u + v, \dots)$ pour f constante vis-à-vis de son premier argument, cela revient au même de considérer $g(u, v, \dots)$ constante vis-à-vis de ses deux premiers arguments et identique à f pour les autres.

$?_1/!_0$ On passe 0 à une fonction linéaire, cela donne donc la fonction nulle.

$?_1/!_1$ On dérive une fonction analytique qui est en fait linéaire, c'est-à-dire $f(x, \dots) = \lambda x g(\dots)$. On a donc $t \times \frac{\partial f}{\partial x}(0, \dots) = t \times \lambda \times g(\dots) = f(t, \dots)$.

$?_1/!_2$ On a $f(x, \dots)$ linéaire en x , vue comme analytique, à laquelle on passe l'argument $u + v$, on calcule donc $f(u + v, \dots) = f(u, \dots) + f(v, \dots)$, dans le premier terme la fonction est constante sur l'argument v et dans le second sur l'argument u , ce qui est représenté par les cellules $?_0$.

$?_2/!_0$ On calcule $g(0, ..) = f(0, 0, \dots)$.

?₂/!₁ On a $g(x, \dots) = f(x, x, \dots)$ et on calcule :

$$t \times \frac{\partial g}{\partial x}(0, \dots) = t \times \frac{\partial f}{\partial x_1}(0, 0, \dots) + t \times \frac{\partial f}{\partial x_2}(0, 0, \dots)$$

où x_1 et la première variable de f et x_2 la seconde.

?₂/!₂ On a $g(u + v, \dots) = f(u + v, u + v, \dots)$ qu'on peut calculer en dupliquant d'abord u et v et en formant après deux arguments $u + v$.

10.2.2 Transfert linéaire de paquets

10.2.2.a Interprétation des réseaux

Un réseau va être interprété comme une soupe de processus communicants. Un lien explicite sera fait ultérieurement avec une traduction d'un calcul de processus dans les réseaux.

Les deux différences avec l'interprétation précédente sont :

- l'absence d'orientation particulière des réseaux,
- et la hiérarchie d'exponentielles, un fil $A - A^\perp$ où A est une variable propositionnelle, représente une connection physique, alors qu'un fil $!A - ?A^\perp$ représente un canal de communication visant à établir une connection $A - A^\perp$ non nécessairement physique.

La modalité exponentielle va donc nous servir à passer du monde physique au monde des communications, qui elles peuvent être routées, terminées, etc. . .

10.2.2.b Interprétation de la somme

La somme est ici interprétée comme le choix non-déterministe, construction usuelle des algèbres de processus. Ainsi le réseau $R + R'$ représente la superposition non-déterministe du réseau R et du réseau R' .

10.2.2.c Interprétation des cellules

Les cellules sont également séparés par dualité : les cellules négatives $\mathfrak{A}, ?^0, ?^1, ?^2$ vont correspondre à des comportements liée à la réception, alors que les cellules positives vont correspondre à l'émission.

Multiplexage Le $\mathfrak{A}$ et le $\otimes$ permettent de grouper des canaux, le premier correspond à l'opération de démultiplexage alors que le second correspond à l'opération de multiplexage. Ces opérations permettent également d'assurer une synchronisation dans les communications.

Requête et paquet La dérélction permet d'effectuer une requête sur un canal, alors que la co-dérélction effectue l'émission d'un paquet

Routage et Co-routage La contraction réalise le routage des différents paquets vers les requêtes, alors que la co-contraction effectue le co-routage des requêtes.

Paquet et requête factices Le co-affaiblissement permet de réaliser un paquet factice, sans contenu réel. C'est à dire un objet vivant exclusivement dans le type des communications. Par dualité, l'affaiblissement permet de réaliser une requête factice.

Levée et rattrapage d'exception L'interprétation précédente peut aussi être vu comme cela : le co-affaiblissement correspond à une erreur d'émission, c'est-à-dire la levée d'une exception, alors que l'affaiblissement correspond au rattrapage d'une telle exception.

10.2.2.d Interprétation des règles

$\mathfrak{V}/\otimes$ il s'agit d'une synchronisation, c'est une propriété élémentaire des opérations de multiplexage et de de-multiplexage.

$?_0/!_0$ Une exception est levée puis rattrapée.

$?_0/!_1$ Un paquet est transmis à une requête factice, on a une erreur de communication.

$?_0/!_2$ On propage la requête factice, pouvant être dupliquée car vivant dans le monde des communications.

$?_2/!_0$ idem, on propage le paquet factice.

$?_1/!_0$ Un paquet factice est reçu, provoquant une erreur.

$?_1/!_1$ On met en communication un paquet et une requête, tout se passe bien et la connection physique, ou tout du moins d'ordre inférieur, peut avoir lieu.

$?_1/!_2$ On effectue le co-routage d'une requête, or la connection d'ordre inférieure derrière la requête ne peut être dupliqué. Il y a deux possibilités, transmettre une requête factice au premier et la requête au second, ou l'inverse.

$?_2/!_1$ idem, on a recours à un paquet factice pour combler les manques.

$?_2/!_2$ Cette règle exprime ici l'indépendance, caractérisée algébriquement par une commutation, des opérations de routage et de co-routage.

10.3 Remarques sur la nature sémantique de la somme

Si on analyse le traitement de la somme fait par les deux interprétations précédentes, on voit que cette somme appartient à la sémantique plus qu'à la syntaxe.

En effet, dans la première interprétation cela est évident, la somme étant celle d'un espace vectoriel sous-jacent. Dans la seconde interprétation, en revanche, c'est plus délicat. Lorsque l'on parle d'un processus comportant des sommes, on est en train de spécifier tous les comportements possibles du processus. Cependant, une exécution se passera dans un terme simple de la somme. De sorte, qu'en exprimant cette somme, on est déjà en train de parler de la sémantique du processus.

Étant donné que, historiquement, les réseaux d'interaction différentiels sont apparus comme un relèvement de la sémantique vers la syntaxe, on peut se demander si la somme a été suffisamment abstraite de la sémantique. Autrement dit, existe-t-il une notion intermédiaire de somme syntaxique dont le quotient, et ainsi la sémantique, serait la somme usuelle ? Posé comme cela, cette question semble insoluble. Cependant, comme nous le verrons dans le chapitre suivant, cette nouvelle notion de somme apparaît naturellement dès que l'on veut regarder les propriétés géométriques, c'est-à-dire les chemins, des réseaux d'interaction différentiels.

Chapitre 11

Géométrie de l'interaction des réseaux d'interaction différentiels

L'essentiel des résultats présentés ici a fait l'objet de la publication [dF08].

Dans ce chapitre on considèrera exclusivement la réduction pertinente des réseaux d'interaction différentiels, c'est-à-dire toutes les règles de réductions hormis les règles pouvant faire apparaître le réseau nul : $?^0/!^1$ et $?^1/!^0$. Ainsi, nous oublierons même de parler du réseau nul dans la mesure où les réseaux non nuls sont stables par la réduction pertinente.

11.1 Chemins dans des sommes

11.1.1 Une approche syntaxique de la somme

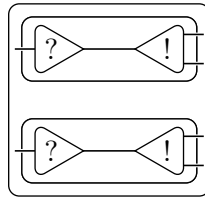
Pour pouvoir parler efficacement de la géométrie des chemins, il va nous falloir distinguer un chemin $\varphi \in \mathfrak{P}(R)$ du même chemin dans une somme $R + R'$, voir même dans une somme $R' + R$. L'approche la plus simple consiste à ne plus considérer la somme comme une opération algébrique mais comme une opération syntaxique. De ce fait nous allons dans ce chapitre abandonner la notion de coefficients dans les sommes.

L'essai apparemment le plus simple serait de définir les réseaux différentiels ainsi :

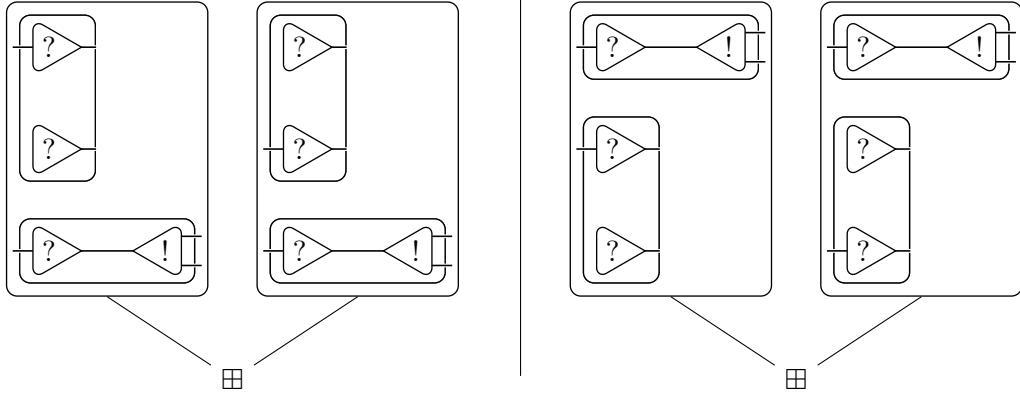
Définition 11.1 *On définit l'ensemble $\mathfrak{R}_\partial$ comme l'ensemble des arbres dont les feuilles sont étiquetées par $\mathfrak{R}\langle P \rangle$ pour un $P \subseteq \mathbb{N}$.*

Plus précisément, $\mathfrak{R}_\partial = \bigcup_{P \subseteq \mathbb{N}} \mathfrak{R}_P$ où $\mathfrak{R}_P$ est le magma librement engendré par $\mathfrak{R}\langle P \rangle$ dont on note l'opération $\boxplus$.

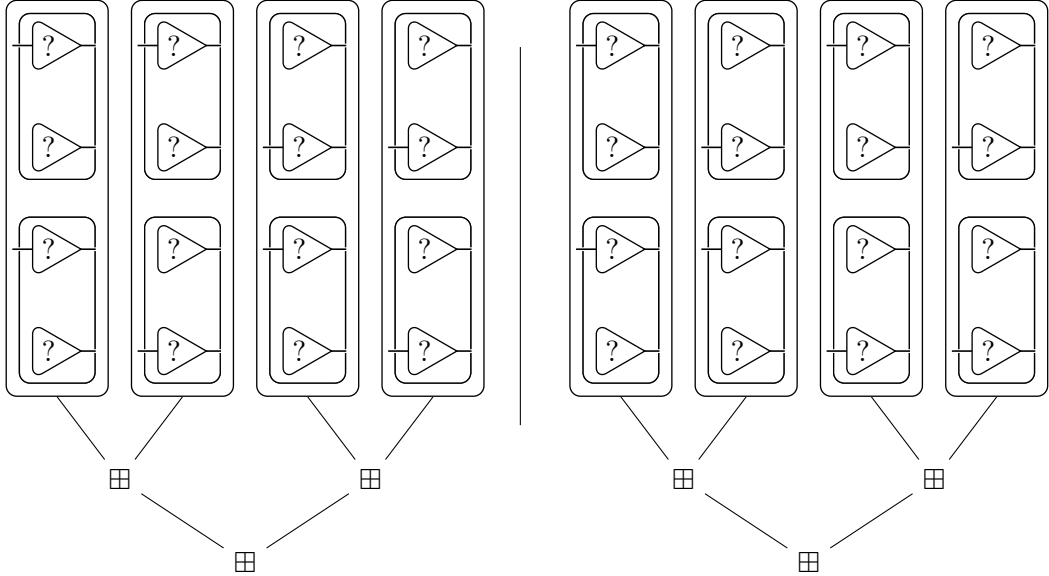
Cependant en faisant ainsi, la construction est trop libre, on perd la confluence forte. En effet, considérons un réseau contenant deux redex $?^1/!^2$:



Nous avons deux choix possibles de réduction, la figure suivante montre les deux résultats :



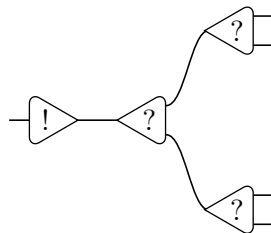
Il est maintenant possible dans les deux cas de réduire le redex restant, donnant ainsi les deux réseaux suivants :



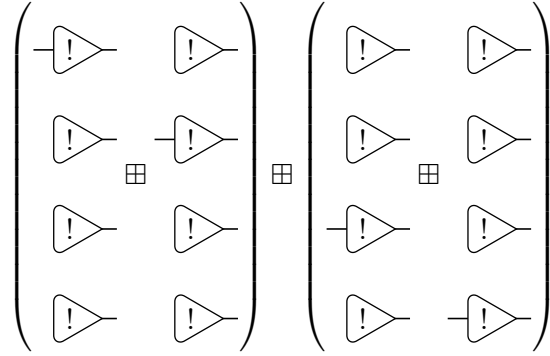
D'une part nous devrions avoir la confluence forte de la réduction, s'exprimant ici par l'égalité de ces deux réseaux, et d'autre part ces deux réseaux sont distincts dans $\mathfrak{R}_\partial$. On aimerait pouvoir quotienter $\mathfrak{R}_\partial$ par la relation $\mathcal{R}_\boxplus$ correspondante :

$$(R_1 \boxplus R_2) \boxplus (R_3 \boxplus R_4) \mathcal{R}_\boxplus (R_1 \boxplus R_3) \boxplus (R_2 \boxplus R_4)$$

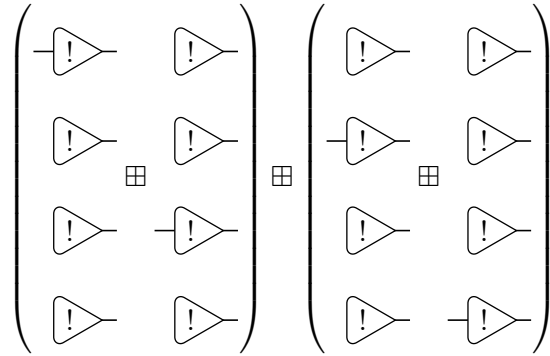
Cependant, en faisant cela on relâche trop la structure. Par exemple, le réseau



se réduit sur la somme



La racine de l'arbre de somme représente alors le premier choix de distribution de la co-déréliction, et nous permet ainsi de faire une dichotomie des chemins en fonction de ce choix. Or ici, le quotient nous ferait complètement perdre cette information en égalant ce réseau avec le réseau suivant :



On voit donc qu'il va falloir restreindre le quotient aux cas posant problème vis-à-vis de la confluence forte, afin de limiter la perte d'informations.

11.1.2 Sommes nommées

Afin de répondre au problème précédent, revenons au sens premier de la somme.

En suivant l'interprétation des réductions du paragraphe 10.2.2, la somme issue de la réduction d'un redex $?^1/!^2$ correspond à un *choix* de distribution de la dérélction, en quelque sorte la somme est indicée par les dérélctions et les co-dérélctions présentes initialement. La relation nécessaire pour récupérer la confluence forte peut alors être vue comme l'indépendance entre deux choix portant sur des ressources distinctes. Les définitions que l'on va donner maintenant tirent parti de cette remarque pour donner une définition juste de la somme.

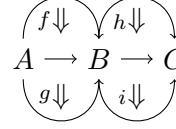
Soit $\mathcal{N}$ un ensemble dénombrable dont les éléments sont appelés *noms*. On considère maintenant que chaque réseau simple est muni d'un étiquetage injectif de ces cellules de dérélction et co-dérélction par $\mathcal{N}$.

Définition 11.2 On note $\mathfrak{R}_\partial = \bigcup_{P \subseteq \mathbb{N}} \mathfrak{R}_P / \mathcal{R}_\boxplus$ où

- $\mathfrak{R}_P$ est librement engendré à partir de $\mathfrak{R}\langle P \rangle$ par les opérations $\boxplus^\alpha$ pour $\alpha \in \mathcal{N}$,
- et, pour $\alpha \neq \beta \in \mathcal{N}$, $(R_1 \boxplus^\beta R_2) \boxplus^\alpha (R_3 \boxplus^\beta R_4) \mathcal{R}_\boxplus (R_1 \boxplus^\alpha R_3) \boxplus^\beta (R_2 \boxplus^\alpha R_4)$

Remarque 11.3 La relation $\mathcal{R}_{\boxplus}$ est plus connu sous le nom de loi d'échange dans les n -

catégories. Par exemple, considérons le diagramme $A \rightarrow B \rightarrow C$, on peut effectuer deux



calculs :

- une composition verticale $A \xrightarrow{fg} B \xrightarrow{hi} C$ suivie d'une composition horizontale $A \xrightarrow{(fg)|| (hi)} C$
- une composition horizontale $A \xrightarrow{f||h} B \xrightarrow{g||i} C$ suivie d'une composition verticale $A \xrightarrow{(f||h)|| (g||i)} C$.

Un des axiomes des n -catégories revient à dire que ces deux résultats sont égaux, on a donc

$$(fg)|| (hi) = (f||h)(g||i)$$

qui est exactement la même relation que celle définissant $\mathcal{R}_{\boxplus}$.

On va maintenant pouvoir décomposer les réductions faisant apparaître des sommes. On note $\xrightarrow{\boxplus}$ la relation $R \xrightarrow{\boxplus} R \boxplus R$ pour tout réseau simple R .

Soient $\mathcal{R}_1$ et $\mathcal{R}_2$ deux règles de réseaux d'interaction, on note $\xrightarrow{\boxplus \mathcal{R}_2}$ la réduction définie par $R \xrightarrow{\boxplus \mathcal{R}_2} R_1 \boxplus R_2$ lorsque $R \xrightarrow{\mathcal{R}_i} R_i$ pour $i \in \{1, 2\}$.

Soient $\mathcal{R}$ une règle de réduction allant de $\mathfrak{R}_\theta^0$ dans $\mathfrak{R}_\theta$, et $R \in \mathfrak{R}_\theta$ comportant une feuille sur laquelle $\mathcal{R}$ s'applique. On peut déduire de R un réseau R' obtenu en appliquant $\mathcal{R}$ sur cette feuille. Soit b la branche menant de cette feuille à la racine, donc composée d'arêtes $\boxplus_g$ ou $\boxplus_d$, on note $\xrightarrow[b]{\mathcal{R}}$ la réduction ainsi définie, s'appliquant sur tous les éléments de $\mathfrak{R}_\theta$ ayant une branche b d'une de leurs feuilles à leur racine. On pose $\xrightarrow{\boxplus} = \cup_b \xrightarrow[b]{\mathcal{R}}$, et de même on pose $\rightarrow^{\boxplus} = \cup_{\mathcal{R}} \xrightarrow{\boxplus}$.

On va montrer que ces définitions permet de récupérer la 2-confluence à partir de la confluence forte feuille à feuille.

Lemme 11.4 Soient $\mathcal{R}_{i,j}$ pour $i, j \in \{1, 2\}$ des règles de réseaux d'interactions, telles que $\mathcal{R}_{i,j}$ et $\mathcal{R}_{i',j'}$ confluent fortement dès que $i \neq i'$.

Si $R \xrightarrow{\mathcal{R}_{1,1} \boxplus \mathcal{R}_{1,2}} R_{1,1} \boxplus R_{1,2}$ et $R \xrightarrow{\mathcal{R}_{2,1} \boxplus \mathcal{R}_{2,2}} R_{2,1} \boxplus R_{2,2}$, avec $\alpha \neq \beta$, alors il existe R' tel que

$$R_{1,1} \boxplus R_{1,2} \xrightarrow[\boxplus_g]{\mathcal{R}_{2,1} \boxplus \mathcal{R}_{2,2}} \xrightarrow[\boxplus_d]{\mathcal{R}_{2,1} \boxplus \mathcal{R}_{2,2}} R'$$

et

$$R_{2,1} \boxplus R_{2,2} \xrightarrow[\boxplus_g]{\mathcal{R}_{1,1} \boxplus \mathcal{R}_{1,2}} \xrightarrow[\boxplus_d]{\mathcal{R}_{1,1} \boxplus \mathcal{R}_{1,2}} R'$$

Preuve On a

$$\begin{array}{ccc} \mathcal{R}_{1,1} \overset{\alpha}{\boxplus} \mathcal{R}_{1,2} & R & \mathcal{R}_{2,1} \overset{\beta}{\boxplus} \mathcal{R}_{2,2} \\ & \swarrow \quad \searrow & \\ R_{1,1} \overset{\alpha}{\boxplus} R_{1,2} & & R_{2,1} \overset{\beta}{\boxplus} R_{2,2} \end{array}$$

et donc par définition $R \xrightarrow{\mathcal{R}_{i,j}} R_{i,j}$, ce qui entraîne l'existence de $R'_{j,j'}$ tels que, pour $j, j' \in \{1, 2\}$,

$$R_{1,j} \xrightarrow{\mathcal{R}_{2,j'}} R'_{j,j'}$$

et

$$R_{2,j'} \xrightarrow{\mathcal{R}_{1,j}} R'_{j,j'}$$

On peut donc refermer le diagramme ainsi :

$$\begin{array}{ccc} \mathcal{R}_{1,1} \overset{\alpha}{\boxplus} \mathcal{R}_{1,2} & R & \mathcal{R}_{2,1} \overset{\beta}{\boxplus} \mathcal{R}_{2,2} \\ & \swarrow \quad \searrow & \\ R_{1,1} \overset{\alpha}{\boxplus} R_{1,2} & & R_{2,1} \overset{\beta}{\boxplus} R_{2,2} \\ \mathcal{R}_{2,1} \overset{\beta}{\boxplus} \mathcal{R}_{2,2} \downarrow \overset{\alpha}{\boxplus}_g & & \mathcal{R}_{1,1} \overset{\alpha}{\boxplus} \mathcal{R}_{1,2} \downarrow \overset{\beta}{\boxplus}_g \\ (R'_{1,1} \overset{\beta}{\boxplus} R'_{1,2}) \overset{\alpha}{\boxplus} R_{1,2} & & (R'_{1,1} \overset{\alpha}{\boxplus} R'_{2,1}) \overset{\beta}{\boxplus} R_{2,2} \\ \mathcal{R}_{2,1} \overset{\beta}{\boxplus} \mathcal{R}_{2,2} \downarrow \overset{\alpha}{\boxplus}_d & & \mathcal{R}_{1,1} \overset{\alpha}{\boxplus} \mathcal{R}_{1,2} \downarrow \overset{\beta}{\boxplus}_d \\ (R'_{1,1} \overset{\beta}{\boxplus} R'_{1,2}) \overset{\alpha}{\boxplus} (R'_{2,1} \overset{\beta}{\boxplus} R'_{2,2}) & \equiv_{\mathcal{R}_{\boxplus}} & (R'_{1,1} \overset{\alpha}{\boxplus} R'_{2,1}) \overset{\beta}{\boxplus} (R'_{1,2} \overset{\alpha}{\boxplus} R'_{2,2}) \end{array}$$

◀

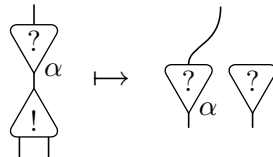
Lemme 11.5 Soient $\mathcal{R}, \mathcal{R}_1$ et $\mathcal{R}_2$ des règles de réseaux d'interactions telles que $\mathcal{R}$ conflue fortement avec $\mathcal{R}_1$ et avec $\mathcal{R}_2$.

Si $R \xrightarrow{\mathcal{R}} R_0$ et $R \xrightarrow{\mathcal{R}_1 \boxplus \mathcal{R}_2} R_1 \overset{\alpha}{\boxplus} R_2$, alors il existe R' tel que $R_0 \xrightarrow{\mathcal{R}_1 \boxplus \mathcal{R}_2}^2 R'$ et $R_1 \overset{\alpha}{\boxplus} R_2 (\xrightarrow{\mathcal{R}})^2 R'$.

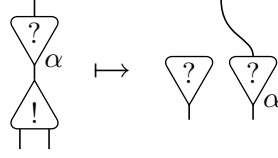
Preuve Étant donné l'hypothèse de confluence forte, on sait qu'il existe R'_1 et R'_2 tels que $R_0 \xrightarrow{\mathcal{R}_i} R'_i$ et $R_i \xrightarrow{\mathcal{R}} R'_i$, pour $i \in \{1, 2\}$.

On vérifie directement que le réseau différentiel $R' = R'_1 \overset{\alpha}{\boxplus} R'_2$ convient. ◀

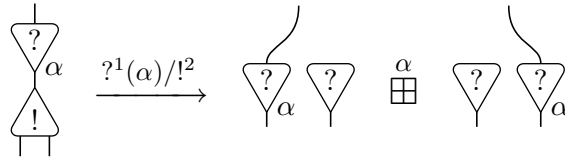
On note $?^1(\alpha)/_g!^2$ (resp. $?^1(\alpha)/_d!^2$) la règle de réduction définie par



(resp. définie par



). On définit maintenant la règle de réduction $?^1(\alpha)/!^2 = ?^1(\alpha)/_g!^2 \overset{\alpha}{\boxplus} ?^1(\alpha)/_d!^2$. Concrètement on a donc :



De manière symétrique on définit $!^1(\alpha)/?^2$. A l'aide des deux lemmes précédents, on peut déduire le théorème suivant :

Théorème 11.6 *La réduction sur $\mathfrak{R}_\partial$ est 2-confluente.*

11.1.3 Chemins

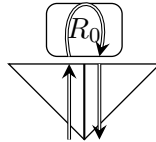
On peut maintenant donner la définition précise d'un chemin dans une somme.

Définition 11.7 *Soit $R \in \mathfrak{R}_\partial$, un chemin dans R est un couple (b, ϕ) où ϕ est un chemin dans un des réseaux simples R_0 présent dans R et b est la branche de R menant de la racine à la feuille R_0 .*

Remarquons que la branche b est en fait une classe d'équivalence modulo $\mathcal{R}_\boxplus$.

On note toujours $\mathfrak{P}(R)$ l'ensemble des chemins de R . Afin de pouvoir utiliser les mêmes définitions associées aux chemins dans ce cadre, nous écrirons un chemin (b, ϕ) comme une succession d'arêtes comme dans le cas usuel, en notant $\overset{\alpha}{\boxplus}_g$ (resp. $\overset{\alpha}{\boxplus}_d$) la branche gauche (resp. droite) d'un $\overset{\alpha}{\boxplus}$ prise du haut vers le bas. Ainsi un chemin (b, ϕ) sera noté $b''\phi b'$ où b' se déduit de b en concaténant les arêtes $\overset{\alpha}{\boxplus}_i$ en partant de la feuille.

Tout se passe donc comme si on considérait le chemin composé ainsi :



où R_0 est le réseau simple dans lequel se situe le *vrai* chemin.

De plus, si pour deux nœuds nommés α et β pour lesquels le quotient s'applique, on a la relation suivante :

$$\overset{\alpha}{\boxplus}_i \overset{\beta}{\boxplus}_j = \overset{\beta}{\boxplus}_j \overset{\alpha}{\boxplus}_i \quad \forall \alpha \neq \beta \in \mathcal{N} \quad \forall i, j \in \{g, d\} \quad (11.1)$$

Une manière directe de s'assurer que l'on puisse réaliser cette égalité, consiste à prendre un élément de $\mathfrak{P}(R)^+$ passant par toutes les branches. Autrement dit, on a, toujours, la relation suivante :

$$\begin{aligned}
 & \begin{array}{c} \alpha^r \beta^r \\ \boxplus_g \boxplus_g \end{array} \varphi_{gg} \begin{array}{c} \beta \alpha \\ \boxplus_g \boxplus_g \end{array} + \begin{array}{c} \alpha^r \beta^r \\ \boxplus_d \boxplus_g \end{array} \varphi_{dg} \begin{array}{c} \beta \alpha \\ \boxplus_g \boxplus_d \end{array} + \begin{array}{c} \alpha^r \beta^r \\ \boxplus_g \boxplus_d \end{array} \varphi_{gd} \begin{array}{c} \beta \alpha \\ \boxplus_d \boxplus_g \end{array} \\
 & + \begin{array}{c} \alpha^r \beta^r \\ \boxplus_d \boxplus_d \end{array} \varphi_{dd} \begin{array}{c} \beta \alpha \\ \boxplus_d \boxplus_d \end{array} \\
 = & \begin{array}{c} \beta^r \alpha^r \\ \boxplus_g \boxplus_g \end{array} \varphi_{gg} \begin{array}{c} \alpha \beta \\ \boxplus_g \boxplus_g \end{array} + \begin{array}{c} \beta^r \alpha^r \\ \boxplus_d \boxplus_g \end{array} \varphi_{gd} \begin{array}{c} \alpha \beta \\ \boxplus_g \boxplus_d \end{array} + \begin{array}{c} \beta^r \alpha^r \\ \boxplus_g \boxplus_d \end{array} \varphi_{dg} \begin{array}{c} \alpha \beta \\ \boxplus_d \boxplus_g \end{array} \\
 & + \begin{array}{c} \beta^r \alpha^r \\ \boxplus_d \boxplus_d \end{array} \varphi_{dd} \begin{array}{c} \alpha \beta \\ \boxplus_d \boxplus_d \end{array}
 \end{aligned} \tag{11.2}$$

11.1.4 Réduction de chemins

Afin de pouvoir définir la réduction de chemin, il va nous falloir d'une part, la définir pour des chemins dans des sommes, et d'autre part, définir l'effet de la création d'une somme par la réduction.

Soit $\varphi \in \mathfrak{P}(R)$, où R est un réseau simple tel que $R \xrightarrow{\mathcal{R}_1 \boxplus \mathcal{R}_2} R_1 \boxplus R_2$ on pose

$$\delta_{\mathcal{R}_1 \boxplus \mathcal{R}_2}(\varphi) = \begin{array}{c} \alpha^r \\ \boxplus_g \end{array} \delta_{\mathcal{R}_1}(\varphi) \begin{array}{c} \alpha \\ \boxplus_g \end{array} + \begin{array}{c} \alpha^r \\ \boxplus_d \end{array} \delta_{\mathcal{R}_2}(\varphi) \begin{array}{c} \alpha \\ \boxplus_d \end{array}$$

Notons que pour les deux règles de cette forme définies précédemment, pour tout chemin traversant le redex on a au moins un des deux termes qui est nul.

Considérons maintenant un réseau d'interaction différentiel $R \xrightarrow[b]{\mathcal{R}} R'$, on va définir la fonction de réduction associée, notée $\delta_{\mathcal{R},b}$. Soit $\varphi \in \mathfrak{P}(R)$, on a donc $\varphi = b'^r \varphi' b'$ avec φ' chemin dans une des feuilles de R . Si $b \neq b'$, le chemin n'est pas affecté par la réduction, et on pose $\delta_{\mathcal{R},b}(\varphi) = \varphi$. Sinon, $b = b'$ et on pose $\delta_{\mathcal{R},b}(\varphi) = b^r \delta_{\mathcal{R}}(\varphi') b$.

Soient $R_g \boxplus R_d \in \mathfrak{R}_\partial$ et $\mathcal{R}$ une règle s'appliquant à la fois sur R_g et sur R_d . Soit φ un chemin de ce réseau, on a soit $\varphi = \begin{array}{c} \alpha^r \\ \boxplus_g \end{array} \varphi_g \begin{array}{c} \alpha \\ \boxplus_g \end{array}$ avec φ_g chemin dans R_g , soit l'inverse. Supposons que l'on soit dans le premier cas, on a alors $\delta_{\mathcal{R}, \begin{array}{c} \alpha \\ \boxplus_d \end{array}}(\varphi) = \varphi$ et ainsi $(\delta_{\mathcal{R}, \begin{array}{c} \alpha \\ \boxplus_g \end{array}} \circ \delta_{\mathcal{R}, \begin{array}{c} \alpha \\ \boxplus_d \end{array}})(\varphi) = \delta_{\mathcal{R}, \begin{array}{c} \alpha \\ \boxplus_g \end{array}}(\varphi)$. Le même raisonnement dans l'autre cas permet de définir la réduction de chemin le long de la double réduction $R_g \boxplus R_d \xrightarrow[\begin{array}{c} \alpha_g \\ \boxplus_g \end{array}]{\mathcal{R}} \xrightarrow[\begin{array}{c} \alpha_d \\ \boxplus_d \end{array}]{\mathcal{R}} R'$. On notera $\delta_{\mathcal{R}, \begin{array}{c} \alpha \\ \boxplus_g + \boxplus_d \end{array}}$ cette application qui admet une expression exacte :

$$\delta_{\mathcal{R}, \begin{array}{c} \alpha \\ \boxplus_g + \boxplus_d \end{array}}(\varphi) = \begin{cases} \begin{array}{c} \alpha^r \\ \boxplus_g \end{array} \delta_{\mathcal{R}}(\varphi') \begin{array}{c} \alpha \\ \boxplus_g \end{array} & \text{si } \varphi = \begin{array}{c} \alpha^r \\ \boxplus_g \end{array} \varphi' \begin{array}{c} \alpha \\ \boxplus_g \end{array} \\ \begin{array}{c} \alpha^r \\ \boxplus_d \end{array} \delta_{\mathcal{R}}(\varphi') \begin{array}{c} \alpha \\ \boxplus_d \end{array} & \text{si } \varphi = \begin{array}{c} \alpha^r \\ \boxplus_d \end{array} \varphi' \begin{array}{c} \alpha \\ \boxplus_d \end{array} \end{cases}$$

11.1.5 Confluence forte de la réduction de chemins

On va énoncer maintenant le pendant du théorème 6.6, notons que l'on a un peu *triché* grâce à la dernière remarque du paragraphe précédent, dans la mesure où l'on compense le fait que la réduction soit seulement 2-confluente par le fait que les réductions de chemins se superposent.

Théorème 11.8 Soient R un réseau d'interaction différentiel sur laquelle s'appliquent deux réductions : $R_1 \xleftarrow[b_1]{\mathcal{R}_1} R \xrightarrow[b_2]{\mathcal{R}_2} R_2$ avec $R_1 \neq R_2$.

Si les deux règles sont à valeur dans $\mathfrak{R}_\partial^0$ ou si l'on a $b_1 \neq b_2$, on a

$$\delta_{\mathcal{R}_1, b_1} \circ \delta_{\mathcal{R}_2, b_2} = \delta_{\mathcal{R}_2, b_2} \circ \delta_{\mathcal{R}_1, b_1}$$

Si $b_1 = b_2 = b$, et que l'une des règles, par exemple $\mathcal{R}_1$, est de la forme $\mathcal{R}_g \boxplus^\alpha \mathcal{R}_d$, l'autre étant à valeur dans $\mathfrak{R}_\partial^0$, on a

$$\delta_{\mathcal{R}_g \boxplus^\alpha \mathcal{R}_d, b} \circ \delta_{\mathcal{R}_2, b} = \delta_{\mathcal{R}_2, b \boxplus_g^\alpha b \boxplus_d^\alpha} \circ \delta_{\mathcal{R}_g \boxplus^\alpha \mathcal{R}_d, b}$$

Si $b_1 = b_2 = b$ et que les deux règles sont de la forme $\mathcal{R}_{i,g} \boxplus^{\alpha_i} \mathcal{R}_{i,d}$, on a

$$\delta_{\mathcal{R}_{1,g} \boxplus^{\alpha_1} \mathcal{R}_{1,d}, b \boxplus_g^{\alpha_2} b \boxplus_d^{\alpha_2}} \circ \delta_{\mathcal{R}_{2,g} \boxplus^{\alpha_2} \mathcal{R}_{2,d}, b} = \delta_{\mathcal{R}_{2,g} \boxplus^{\alpha_2} \mathcal{R}_{2,d}, b \boxplus_g^{\alpha_1} b \boxplus_d^{\alpha_1}} \circ \delta_{\mathcal{R}_{1,g} \boxplus^{\alpha_1} \mathcal{R}_{1,d}, b}$$

Preuve Le théorème énonce en fait trois propriétés distinctes que l'on va prouver à part.

1. Si $b_1 \neq b_2$ on a pour tout chemin φ dans R , soit $\delta_{\mathcal{R}_1, b_1}(\varphi) = \varphi$, soit $\delta_{\mathcal{R}_2, b_2}(\varphi) = \varphi$, soit les deux. Si $b_1 = b_2$ et que les deux règles sont à valeurs dans $\mathfrak{R}_\partial^0$, on est ramené au théorème 6.6.
2. Supposons maintenant que $b_1 = b_2 = b$, $\mathcal{R}_1 = \mathcal{R}_g \boxplus^\alpha \mathcal{R}_d$ et que $\mathcal{R}_2$ est à valeur dans $\mathfrak{R}_\partial^0$. Soit φ chemin dans R , le résultat est évident si φ n'est pas le long de b , posons ainsi $\varphi = b^r \varphi' b$. On a

$$\begin{aligned} & (\delta_{\mathcal{R}_g \boxplus^\alpha \mathcal{R}_d, b} \circ \delta_{\mathcal{R}_2, b})(\varphi) \\ &= b^r (\boxplus_g^{\alpha r} (\delta_{\mathcal{R}_g} \circ \delta_{\mathcal{R}_2})(\varphi') \boxplus_g^\alpha + \boxplus_d^{\alpha r} (\delta_{\mathcal{R}_d} \circ \delta_{\mathcal{R}_2})(\varphi') \boxplus_d^\alpha) b \end{aligned}$$

et

$$\begin{aligned} & (\delta_{\mathcal{R}_2, b \boxplus_g^\alpha b \boxplus_d^\alpha} \circ \delta_{\mathcal{R}_g \boxplus^\alpha \mathcal{R}_d, b})(\varphi) \\ &= \delta_{\mathcal{R}_2, b \boxplus_g^\alpha b \boxplus_d^\alpha} (b^r (\boxplus_g^{\alpha r} \delta_{\mathcal{R}_g}(\varphi') \boxplus_g^\alpha + \boxplus_d^{\alpha r} \delta_{\mathcal{R}_d}(\varphi') \boxplus_d^\alpha) b) \\ &= b^r (\boxplus_g^{\alpha r} (\delta_{\mathcal{R}_2} \circ \delta_{\mathcal{R}_g})(\varphi') \boxplus_g^\alpha + \boxplus_d^{\alpha r} (\delta_{\mathcal{R}_2} \circ \delta_{\mathcal{R}_d})(\varphi') \boxplus_d^\alpha) b \end{aligned}$$

Et on conclut en appliquant le théorème 6.6.

3. Supposons que $b_1 = b_2 = b$ et que $\mathcal{R}_i = \mathcal{R}_{i,g} \boxplus^{\alpha_i} \mathcal{R}_{i,d}$, pour $i \in \{1, 2\}$. Soit φ chemin dans R , il suffit là encore de ne considérer que le cas où $\varphi = b^r \varphi' b$. On a

$$\begin{aligned} & (\delta_{\mathcal{R}_{1,g} \boxplus^{\alpha_1} \mathcal{R}_{1,d}, b \boxplus_g^{\alpha_2} b \boxplus_d^{\alpha_2}} \circ \delta_{\mathcal{R}_{2,g} \boxplus^{\alpha_2} \mathcal{R}_{2,d}, b})(\varphi) \\ &= b^r \sum_{k, l \in \{g, d\}} (\boxplus_k^{\alpha_2 r} \boxplus_l^{\alpha_1 r} (\delta_{\mathcal{R}_{1,l}} \circ \delta_{\mathcal{R}_{2,k}})(\varphi') \boxplus_l^{\alpha_1} \boxplus_k^{\alpha_2}) b \end{aligned}$$

et

$$\begin{aligned} & (\delta_{\mathcal{R}_{2,g} \boxplus^{\alpha_2} \mathcal{R}_{2,d}, b \boxplus^{\alpha_1} g + b \boxplus^{\alpha_1} d} \circ \delta_{\mathcal{R}_{1,g} \boxplus^{\alpha_1} \mathcal{R}_{1,d}, b})(\varphi) \\ &= b^r \sum_{k,l \in \{g,d\}} (\boxplus_l^{\alpha_1 r} \boxplus_k^{\alpha_2 r} (\delta_{\mathcal{R}_{2,k}} \circ \delta_{\mathcal{R}_{1,l}})(\varphi')) \boxplus_k^{\alpha_2} \boxplus_l^{\alpha_1} b \end{aligned}$$

On conclut en appliquant le théorème 6.6 puis l'équation (11.2).

✎

Corollaire 11.9 *La notion de persistance fait sens pour les chemins dans des réseaux d'interaction différentiels.*

11.2 L'algèbre ∂L^*

11.2.1 Définition

Soit

$$\partial L^* = \langle p, q, r_!, s_!, r_?, s_?, d_{?,\alpha}, d_{!,\alpha}, u_\alpha, v_\alpha, e_\alpha, \dots \mid r \rangle_0^1$$

où les points de suspension représentent les générateurs pour d'autre choix de nom pris dans $\mathcal{N}$, et r représente les relations suivantes (avec les notations du paragraphe 7.3.1) :

$$p \quad \perp \quad q \tag{11.3}$$

$$r_t \quad \perp \quad s_t \tag{11.4}$$

$$u_\alpha \quad \perp^+ \quad v_\alpha \tag{11.5}$$

$$e_\alpha e_\alpha = e_\alpha \tag{11.6}$$

$$\{r_!, s_!\} \leftrightarrow \{r_?^*, s_?^*\} \tag{11.7}$$

$$\begin{aligned} \{u_\alpha, v_\alpha, e_\alpha\} &\stackrel{*}{\leftrightarrow} \{p, q, r_?, s_?, r_!, s_!\} \cup \\ &\bigcup_{\beta \in \mathcal{N} - \{\alpha\}} \{d_{\beta,?}, d_{\beta,!}, u_\beta, v_\beta, e_\beta\} \end{aligned} \tag{11.8}$$

$$d_{\alpha,t}^* r_{t'} = u_\alpha d_{\alpha,t}^* u_\alpha^* \text{ avec } t \neq t' \in \{!, ?\} \tag{11.9}$$

$$d_{\alpha,t}^* s_{t'} = v_\alpha d_{\alpha,t}^* v_\alpha^* \text{ avec } t \neq t' \in \{!, ?\} \tag{11.10}$$

$$d_{\alpha,!}^* d_{\beta,?} = e_\alpha e_\beta \tag{11.11}$$

Les notations étant celles du paragraphe 7.3.1. ∂L^* est normal, mais il ne vérifie pas la propriété ab^* . En effet, un élément comme $p^* r_?$ n'admet pas d'écriture sous cette forme.

On va distinguer deux sous-monoïdes de ∂L^* qui nous permettront de caractériser la forme des poids que l'on va associer aux chemins d'un réseau différentiel. On appelle ∂L_{me}^{*+} le monoïde généré par

$$p, q, r_?, s_?, r_!, s_!, d_{\alpha,?}, d_{\alpha,!}, \dots$$

et ∂L_a^{*+} le monoïde généré par

$$u_\alpha, v_\alpha, e_\alpha, \dots$$

11.2.2 Facteur de normalisation

Définition 11.10 Soit R un réseau différentiel, on note $\mathcal{N}(R)$ l'ensemble des noms présents dans R . De même, on note $\mathcal{N}(\varphi)$ l'ensemble des noms de cellules ou de nœud traversés par le chemin φ . Soit $w \in \partial L^*$, on note $\mathcal{N}(w)$ l'intersection pour toute représentation de w en termes des générateurs dans ∂L^* de l'ensemble des noms apparaissant en étiquette des générateurs contenu dans la représentation.

Remarque 11.11 Les ensembles de noms ne sont modifiés au cours de la réduction que par les règles suivantes :

- $?^1(\alpha)/!^1(\beta)$: les noms α et β disparaissent
- $?^1(\alpha)/!^2$: α devient un nom associé à tous les réduits des chemins présents dans le réseau simple où se situe la réduction. En effet, ils doivent maintenant passer par une branche $\overset{\alpha}{\boxplus} \square$.
- $!^1(\alpha)/?^2$: idem

Définition 11.12 Soient φ et φ' deux chemins dans des réseaux différentiels, éventuellement distincts, on note

$$n_{\varphi'}(\varphi) = \prod_{\alpha \in \mathcal{N}(\varphi') - \mathcal{N}(\varphi)} e_{\alpha}$$

On appelle cet élément de ∂L^* le facteur de normalisation de φ par rapport à φ' . On définit de même $n_x(y)$ dès que x et y ont des ensembles de noms associés par la définition précédente.

L'intérêt de ces facteurs vient du fait qu'ils permettent de comparer des poids de chemins lorsqu'ils n'utilisent pas les mêmes noms. On a évidemment $\mathcal{N}(\mathbf{w}(\varphi)n_{\varphi'}(\varphi)) = \mathcal{N}(\mathbf{w}(\varphi')n_{\varphi}(\varphi'))$.

Les deux lemmes suivants permettent d'analyser complètement l'évolution des facteurs le long de la réduction des chemins.

Lemme 11.13 Soient $R \in \mathfrak{R}_{\partial}$ tel que $R \xrightarrow{\mathcal{R}} R'$ et $\varphi \in \mathfrak{P}(R)$, avec $\delta_{\mathcal{R}}(\varphi) = \varphi'$, on a

1. $\mathcal{N}(\varphi') \subseteq \mathcal{N}(\varphi)$
2. $n_R(\varphi') = n_R(\varphi)n_{\varphi}(\varphi')$
3. si $\mathcal{R} = ?^1(\alpha)/!^1(\beta)$ et φ traverse le redex, on a $n_{\varphi}(\varphi') = e_{\alpha}e_{\beta}$. Dans tous les autres cas, on a $n_{\varphi}(\varphi') = 1$.

Preuve En vertu de la remarque 11.11, il suffit de ne considérer que trois règles, dont deux sont duales l'une de l'autre. L'inclusion est triviale, le seul cas où des noms peuvent être ajoutés étant exclu par le fait que φ se réduit en un seul chemin φ' . L'égalité est alors conséquence directe de cette inclusion. La troisième propriété découle de l'observation directe des règles de réduction. ◀

Lemme 11.14 Soient R un réseau d'interaction différentiel et $b^r \varphi b \in \mathfrak{P}(R)$, tels que $\alpha \notin \mathcal{N}(b)$,

$$R \xrightarrow[b]{?^1(\alpha)/!^2} R'$$

et

$$\delta_{?^1(\alpha)/!^2, b}(b^r \varphi b) = b^r \overset{\alpha^r}{\boxplus}_g \varphi \overset{\alpha}{\boxplus}_g b + b^r \overset{\alpha^r}{\boxplus}_d \varphi \overset{\alpha}{\boxplus}_d b = \phi_1 + \phi_2$$

On a alors $n_{\phi_i}(b^r \varphi b) = e_{\alpha}$ et $n_{b^r \varphi b}(\phi_i) = 1$ pour $i \in \{1, 2\}$. De manière duale, on a la même propriété pour $!^1(\alpha)/?^2$.

Preuve Là encore, conformément à la remarque 11.11, il s'agit de l'unique cas où l'on rajoute un nom. On est assuré que α n'apparaît pas dans φ car sinon il traverserait le redex et ne pourrait être réduit sur deux chemins. Or, α apparaît trivialement dans les ϕ_i par l'intermédiaire des arêtes $\overset{\alpha}{\boxplus}_{\square}$. $\blacktriangleright$

11.2.3 Commutants

Définition 11.15 Soit $w \in \partial L^*$, on appelle commutant de w l'ensemble $w^c \in \partial L^*$ défini par $w^c = \{w' \in \partial L^* \mid ww' = w'w\}$.

Fait 11.16 Soit $x \in \{e_\alpha, u_\alpha, v_\alpha, d_{\alpha,?}, d_{\alpha,!}\}$, on a $x^c = \{x\} \cup \{w \in \partial L^* \mid \alpha \notin \mathcal{N}(w)\}$.

Preuve C'est une application immédiate de la relation 11.8. $\blacktriangleright$

11.3 Pondération des réseaux d'interaction différentiels

11.3.1 Définition

On va définir une pondération des réseaux simples à valeur dans ∂L^* . Dans la mesure où certaines cellules sont nommées, il ne va pas être possible de définir une pondération générique au sens précédent. On écrit $w(t_\alpha, i)$ pour le poids de la i -ème arête d'une cellule de symbole t et de nom α . Ainsi, on pose

$$\begin{aligned} w(\mathfrak{A}, 1) &= p & w(\mathfrak{A}, 2) &= q & w(\otimes, 1) &= q & w(\otimes, 2) &= p \\ w(?^2, 1) &= r? & w(?^2, 2) &= s? & w(!^2, 1) &= r! & w(!^2, 2) &= s! \\ w(?^1_\alpha, 1) &= d_{\alpha,?} & w(!^1_\alpha, 1) &= d_{\alpha,!} \end{aligned}$$

On étend la pondération aux réseaux différentiels en posant

$$w(\overset{\alpha}{\boxplus}_g) = u_\alpha \quad w(\overset{\alpha}{\boxplus}_d) = v_\alpha$$

Etant donné que tout chemin φ se décompose sous la forme $b^r \varphi_0 b$ où b est une branche de l'arbre de somme, on note $w_{\boxplus}(\varphi) = w(b)$ et $w_s(\varphi) = w(\varphi_0)$. De sorte que, $w(\varphi) = w_{\boxplus}(b)w_s(\varphi_0)w_{\boxplus}(b)^*$.

11.3.2 Fidélité

Le lemme suivant est celui sur lequel vont reposer tous les résultats sur cette pondération.

Lemme 11.17 (Lemme fondamental) Soit R un réseau simple tel que $R \xrightarrow{\mathcal{R}} R'$, et $\varphi \in \mathfrak{P}_f(R)$ tel que φ traverse le redex correspondant à la réduction, on a

- soit $\delta_{\mathcal{R}}(\varphi) = \varphi'$ et $w(\varphi) = n_\varphi(\varphi')w(\varphi')$
- soit $\delta_{\mathcal{R}}(\varphi) = 0$ et $w(\varphi) = 0$.

Preuve Seuls deux cas sont importants, les autres relevant d'une vérification directe des équation 7.1.

Notons c_1 et c_2 les deux cellules réduites, on peut décomposer φ selon ses passages à travers le redex. Posons $\varphi = \varphi_0(\prod_i \rho_i \varphi_i)$ avec φ_j ne traversant pas le redex et ρ_i de la forme $c_{1,j}c_{2,k}^r$ ou $c_{2,k}c_{1,j}^r$.

Supposons que $\mathcal{R} = ?^1(\alpha)/!^1(\beta)$, on a alors

$$\mathbf{w}(\varphi) = \left(\prod_i \mathbf{w}(\varphi_i) w_i \right) \mathbf{w}(\varphi_0)$$

avec $w_i = d_{\alpha,?}^* d_{\beta,!}$ ou $d_{\beta,!}^* d_{\alpha,?}$, c'est-à-dire dans tous les cas $w_i = e_\alpha e_\beta$. Or, par construction $\alpha, \beta \notin \mathcal{N}(\varphi_i)$, et le fait 11.16 s'applique. On a donc $\mathbf{w}(\varphi) = (\prod_i \mathbf{w}(\varphi_i) \mathbf{w}(\varphi_0)) (\prod_i e_\alpha) (\prod_i e_\beta) = \mathbf{w}(\varphi') e_\alpha e_\beta = \mathbf{w}(\varphi') n_\varphi(\varphi')$.

Supposons que $\mathcal{R} = ?^1(\alpha)/!^2$, on a alors

$$\mathbf{w}(\varphi) = \left(\prod_i \mathbf{w}(\varphi_i) w_i \right) \mathbf{w}(\varphi_0)$$

avec w_i de la forme $d_{\alpha,?}^* r_!$, $d_{\alpha,?}^* s_! r_!^* d_{\alpha,?}$ ou $r_?^* d_{\alpha,?}$. Après simplification, w_i est donc de la forme $u_\alpha d_{\alpha,?}^\square u_\alpha^*$ ou $v_\alpha d_{\alpha,?}^\square v_\alpha^*$, avec $\square$ égal à rien ou $\star$.

On va se contenter de deux passages successifs de φ à travers le redex, toute la difficulté résidant ici. On a alors $\mathbf{w}(\varphi) = \mathbf{w}(\varphi_2) w_2 \mathbf{w}(\varphi_1) w_1 \mathbf{w}(\varphi_0)$. Pour continuer, nous allons regarder deux cas de figures. Premièrement, supposons que φ passe toujours du même côté de la cellule de co-contraction, par exemple par le premier port auxiliaire. Par la réduction a donc $\delta(\varphi) = \varphi'$ avec φ' commençant par $\overset{\alpha}{\boxplus}_g^r$ et finissant par $\overset{\alpha}{\boxplus}_g$. On a ainsi $w_i = u_\alpha d_{\alpha,?}^{t_i} u_\alpha^*$ avec t_i étant rien ou $\star$. Donc, $\mathbf{w}(\varphi) = \mathbf{w}(\varphi_2) u_\alpha d_{\alpha,?}^{t_2} u_\alpha^* \mathbf{w}(\varphi_1) u_\alpha d_{\alpha,?}^{t_1} u_\alpha^* \mathbf{w}(\varphi_0) = u_\alpha \mathbf{w}(\varphi_2) d_{\alpha,?}^{t_2} \mathbf{w}(\varphi_1) d_{\alpha,?}^{t_1} \mathbf{w}(\varphi_0) u_\alpha^* = \mathbf{w}(\varphi')$, en utilisant la relation 11.5 et le fait 11.16. Notons que l'on peut effectivement appliquer les relations de commutations car α n'apparaît pas dans les φ_i . En effet, dans la mesure où l'on est dans un réseau simple, l'étiquetage des cellules est injectif et on ne peut pas avoir de cellule nommée α en dehors du redex.

Le second cas correspond à φ passant des deux côtés de la cellule de co-contraction. Par exemple, $w_1 = u_\alpha d_{\alpha,?}^{t_1} u_\alpha^*$ et $w_2 = v_\alpha d_{\alpha,?}^{t_2} v_\alpha^*$. On a alors $\mathbf{w}(\varphi) = \mathbf{w}(\varphi_2) u_\alpha d_{\alpha,?}^{t_2} u_\alpha^* \mathbf{w}(\varphi_1) v_\alpha d_{\alpha,?}^{t_1} v_\alpha^* \mathbf{w}(\varphi_0) = w u_\alpha^* v_\alpha w' = 0$. Or, dans ce cas-là effectivement $\delta(\varphi) = 0$.

La figure 11.1 présente graphiquement ces deux cas.

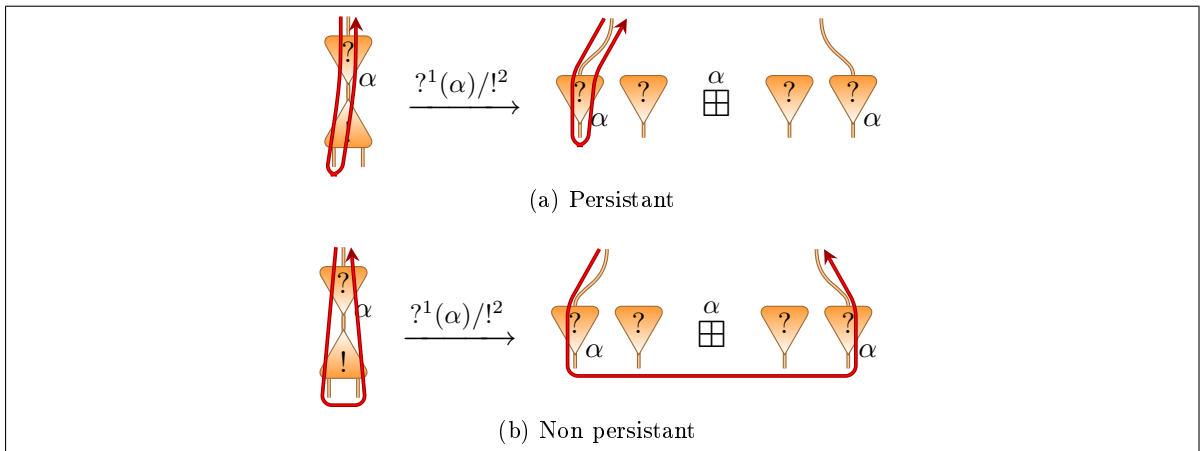


Figure 11.1: Représentation des deux cas apparaissant dans la preuve du lemme 11.17 pour $\mathcal{R} = ?^1(\alpha)/!^2$

Remarque 11.18 Cette preuve fait apparaître très clairement la spécificité de la géométrie des réseaux différentiels. Alors que dans les logiques présentées jusqu'ici la concaténée de deux chemins persistants assez longs est toujours persistante, nous avons ici un phénomène d'interférence lié au non-déterminisme. En effet la concaténée de deux chemins se situant après réduction dans des termes différents d'une somme est non persistante. C'est d'ailleurs la raison principale qui empêche de donner une preuve par simple vérification des relations 7.1.

Il est à noter qu'un cas apparemment similaire apparaît pour un chemin pal/pal dupliqué et passant de deux côtés distincts de la cellule de duplication, ou un chemin se réduisant vers un chemin pal/pal, plus communément appelé un $!$ -cycle. Cependant, de tels chemins ne peuvent être divisés en sous-chemins assez long, ce qui ne contredit donc pas l'affirmation précédente.

Le terme d'interférence pour décrire ce phénomène n'est pas innocent, en effet, un chemin assez long peut être vu comme un rayon de lumière, le terme de la somme dans lequel il se situe correspondrait alors à une sorte de polarisation du signal. On voit alors que la superposition de deux signaux de polarités orthogonales s'annule, ce qui est appelé en optique ondulatoire : l'interférence. Cette analogie physique est d'autant plus profonde que nos poids sont des opérateurs sur un espace de Hilbert.

Il reste un seul cas de réduction de chemin non couvert par le lemme, il est résolu par le lemme suivant.

Lemme 11.19 Soient R un réseau simple tel que $R \xrightarrow{?^1(\alpha)/!^2} R_1 \boxplus R_2$, et $\varphi \in \mathfrak{P}(R)$ tel que $\alpha \notin \mathcal{N}(\varphi)$. On a alors $\delta(\varphi) = \boxplus_g^{\alpha^r} \varphi \boxplus_g^{\alpha} + \boxplus_d^{\alpha^r} \varphi \boxplus_d^{\alpha}$ et $\mathbf{w}(\varphi) = \mathbf{w}(\delta(\varphi))$.

Preuve Remarquons tout d'abord que $\delta(\varphi)$ est donné sous cette forme par les définitions de la réduction de chemin. On a donc

$$\begin{aligned} \mathbf{w}(\boxplus_g^{\alpha^r} \varphi \boxplus_g^{\alpha} + \boxplus_d^{\alpha^r} \varphi \boxplus_d^{\alpha}) &= u_{\alpha} \mathbf{w}(\varphi) u_{\alpha}^* + v_{\alpha} \mathbf{w}(\varphi) v_{\alpha}^* \\ &= (u_{\alpha} u_{\alpha}^* + v_{\alpha} v_{\alpha}^*) \mathbf{w}(\varphi) = \mathbf{w}(\varphi) \end{aligned}$$

◀

On peut résumer les deux lemmes précédents par le théorème suivant.

Théorème 11.20 Soit R un réseau simple et $R \rightarrow R'$ une réduction, on a

$$\forall \varphi \in \mathfrak{P}_f(R), \mathbf{w}(\varphi) = n_{\mathbf{w}(\varphi)}(\mathbf{w}(\delta(\varphi))) \mathbf{w}(\delta(\varphi))$$

Autrement dit, $\mathbf{w}$ est fortement fidèle à un facteur de normalisation près.

Remarquons que l'utilisation du facteur $n_{\mathbf{w}(\varphi)}(\mathbf{w}(\delta(\varphi)))$ à la place du facteur $n_{\varphi}(\delta(\varphi))$ est due au lemme 11.19.

On déduit du théorème précédent un résultat concernant les réseaux différentiels : il suffit de normaliser après être passé à travers la branche menant au réseau simple réduit.

Corollaire 11.21 Soit $R \xrightarrow[b]{\mathcal{R}} R'$ et $\varphi = b^r \varphi' b \in \mathfrak{P}_f(R)$, on a

$$\mathbf{w}(\delta_{\mathcal{R},b}(\varphi')) = \mathbf{w}(b) n \mathbf{w}(\varphi') \mathbf{w}(b)^* = \mathbf{w}(b) n \mathbf{w}(\varphi') n^* \mathbf{w}(b)^*$$

où

$$n = n_{\mathbf{w}(\varphi')}(\mathbf{w}[\delta_{\mathcal{R}}(\varphi')]) = \begin{cases} e_{\alpha} e_{\beta} & \text{si } \mathcal{R} = ?^1(\alpha)/!^1(\beta) \text{ et } \alpha, \beta \in \mathcal{N}(\varphi') \\ 1 & \text{sinon} \end{cases}$$

11.3.3 Structure des poids

On peut adapter la proposition 7.8 valide pour les réseaux d'interaction standards au cas des réseaux différentiels. Tout d'abord le lemme suivant donne la forme générale d'un réseau différentiel normal.

Lemme 11.22 *Soit R un réseau différentiel bien typé. Si R est normal alors R est un arbre de somme dont les feuilles sont des réseaux simples où les fils actifs lient au moins une cellule d'affaiblissement ou de co-affaiblissement.*

Preuve Preuve immédiate compte tenu de la limitation à la réduction faible des réseaux différentiels. ◀

Corollaire 11.23 *Les poids de chemins dans un réseau différentiel bien typé et normal sont de la forme $\alpha ab^* \alpha^*$ avec $\alpha \in \partial L_a^{*+}$ et $a, b \in \partial L_{me}^{*+}$.*

Théorème 11.24 (Forme cab^*c^*) *Soient R un réseau différentiel bien typé et normalisable, et $\varphi \in \mathfrak{P}_f(R)$. On a $w(\varphi) = 0$ ou $w(\varphi) = cab^*c^*$, avec $c \in \partial L_a^{*+}$ et $a, b \in \partial L_{me}^{*+}$.*

On note alors $t_\varphi = c$ et on l'appelle la tranche de φ . On a $t_\varphi = c_{\alpha_1} \dots c_{\alpha_n}$ avec $c_\alpha = d_1 \dots d_m e$, où $d_i = u_\alpha$ ou v_α et $e = 1$ ou e_α . On note $\mathsf{T} = \{t_\varphi \mid R \text{ réseau et } \varphi \in \mathfrak{P}_f(R)\}$.

Preuve Tout d'abord, remarquons que les poids de chemins dans un réseau normal ont cette forme.

Supposons maintenant que $R \rightarrow R'$ et que les poids non nuls de chemins de R' sont de la forme désiré. Soit $\varphi \in \mathfrak{P}_f(R)$ tel que $w(\varphi) \neq 0$. En vertu du corollaire 11.21 on a donc $w(\varphi) = c'ncab^*c^*n^*c'^*$ où n est de la forme 1 ou $e_\alpha e_\beta$ et $c' \in \partial L_a^{*+}$. On peut donc poser $c'' = c'nc$ et on obtient $w(\varphi) = c''ab^*c''^*$. ◀

11.3.4 Pondération saturée

A l'aide du théorème précédent, il est possible de saturer les poids en rajoutant tous les facteurs de normalisation possibles. Ainsi, on récupère une pondération fortement fidèle.

Définition-Propriété 11.25 *Soit R un réseau différentiel et $\varphi \in \mathfrak{P}_f(R)$ régulier avec $w(\varphi) = cab^*c^*$, on pose $\bar{w}(\varphi) = cnab^*n^*c^*$ où $n = \prod_{\alpha \in \mathcal{N}(c) - \mathcal{N}(ab^*)} e_\alpha$.*

$\bar{w}$ est une pondération fortement fidèle.

Corollaire 11.26 *La quantité $\text{NEx}(R) = \sum_{\varphi \in \mathfrak{P}_f(R)} \bar{w}(\varphi)$ est invariante par réduction.*

11.3.5 e_α

Si il est assez raisonnable d'avoir introduit la plupart des générateurs et des relations de ∂L^* pour représenter fidèlement le calcul, une partie de ceux-ci est assez étrange. Pourquoi ne pas avoir remplacer l'équation (11.11) par l'équation beaucoup plus simple

$$d_{\alpha,?}^* d_{\beta,!} = 1 \tag{11.12}$$

et, ainsi supprimer toutes références aux générateurs e_α ? En effet, ils n'ont pas d'existence réelle dans la définition de la pondération, et ils n'apparaissent que pour combler un manque artificiel.

En fait on va voir qu'ils sont absolument nécessaire et que toute la théorie s'écroule si on les supprime.

Considérons donc le monoïde inversif avec zéro ∂L^* obtenu en substituant l'équation (11.11) par l'équation (11.12) et en supprimant les générateurs e_α . La définition de la pondération dans ∂L^* ne faisant pas référence aux e_α elle s'adapte directement pour donner une pondération w à valeur dans ∂L^* .

L'inverse étant unique on a $d_{\alpha,?}^* = d_{\beta,!}^*$ et ce quel que soient α et β . On en déduit donc $\forall \alpha, \beta$ que $d_{\alpha,?} = d_{\beta,!} = d$. En utilisant l'équation (11.9) on en déduit que pour $\alpha \neq \beta$

$$u_\alpha du_\alpha^* = u_\beta du_\beta^*$$

et donc en multipliant par v_α à droite on obtient

$$0 = u_\beta dv_\alpha u_\beta^*$$

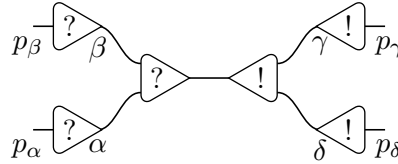
et en multipliant par $v_\alpha^* d^* u_\beta^*$ à gauche et par u_β à droite on obtient $0 = 1$.

Ainsi $\partial L^* = 0$.

Remarque 11.27 *En fait derrière le procès des e_α se cachait le procès des noms. Ceux-ci ont été justifiés pour la définition des chemins, cependant, rien n'assurait qu'ils soient indispensables. On a ici la confirmation par la géométrie des réseaux différentiels que les noms sont nécessaires.*

11.3.6 Exemple

Soit R le réseau simple suivant



On a $\mathfrak{P}_f(R) = \{\varphi_{\alpha\gamma}, \varphi_{\alpha\gamma}^r, \varphi_{\alpha\delta}, \varphi_{\alpha\delta}^r, \varphi_{\beta\gamma}, \varphi_{\beta\gamma}^r, \varphi_{\beta\delta}, \varphi_{\beta\delta}^r\}$ où φ_{nm} est l'unique chemin allant de p_n à p_m . Calculons les poids associés à ces chemins :

$$\begin{aligned} w(\varphi_{\alpha\gamma}) &= d_{\gamma,!}^* r_{?}^* r_{?}^* d_{\alpha,?} = d_{\gamma,!}^* r_{?}^* r_{?}^* d_{\alpha,?} = u_\gamma d_{\gamma,!}^* u_\gamma^* u_\alpha d_{\alpha,?} u_\alpha^* \\ &= u_\gamma u_\alpha d_{\gamma,!}^* d_{\alpha,?} u_\gamma^* u_\alpha^* = u_\alpha e_\alpha u_\alpha^* u_\gamma e_\gamma u_\gamma^* = \overline{w}(\varphi_{\alpha\gamma}) \end{aligned}$$

$$w(\varphi_{\alpha\delta}) = v_\alpha e_\alpha v_\alpha^* u_\delta e_\delta u_\delta^* = \overline{w}(\varphi_{\alpha\delta})$$

$$w(\varphi_{\beta\gamma}) = u_\beta e_\beta u_\beta^* v_\gamma e_\gamma v_\gamma^* = \overline{w}(\varphi_{\beta\gamma})$$

$$w(\varphi_{\beta\delta}) = v_\beta e_\beta v_\beta^* v_\delta e_\delta v_\delta^* = \overline{w}(\varphi_{\beta\delta})$$

Après réduction on obtient le réseau différentiel donné à la figure 11.2 où l'on a volontairement laissé les redex $?^1/!^1$ pour visualiser les noms. Les deux réseaux simples encadrés sont les deux seuls qui ne se réduisent pas en 0 après réduction complète. On a également représenté les différentes réductions des chemins précédents. Calculons les poids des réduits :

$$w(\delta(\varphi_{\alpha\gamma})) = u_\alpha u_\gamma u_\gamma^* u_\alpha^*$$

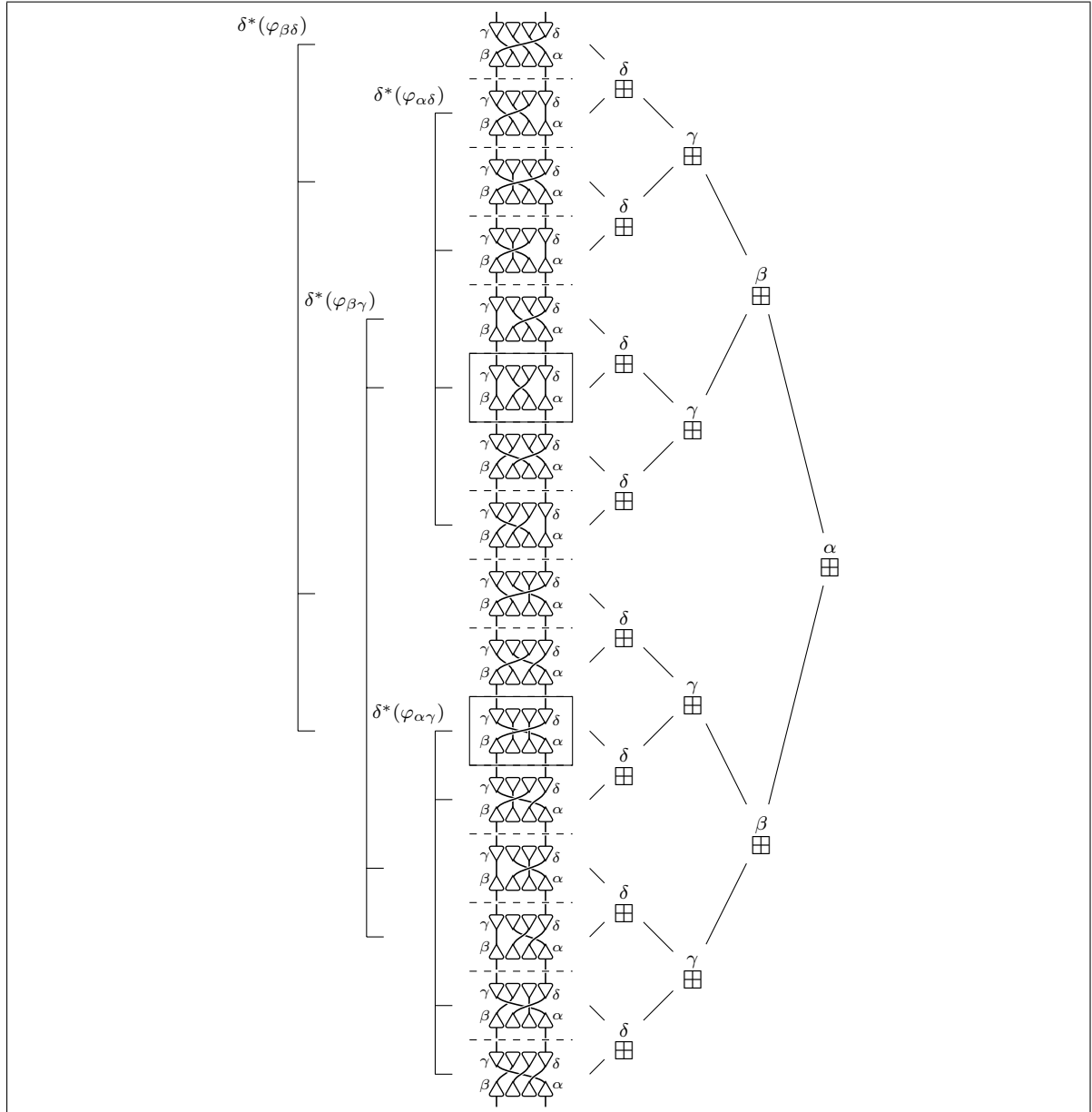


Figure 11.2: Réduit du réseau donné au paragraphe 11.3.6

$$\mathbf{w}(\delta(\varphi_{\alpha\delta})) = v_{\alpha}u_{\delta}u_{\delta}^*v_{\alpha}^*$$

$$\mathbf{w}(\delta(\varphi_{\beta\gamma})) = u_{\beta}v_{\gamma}v_{\gamma}^*u_{\beta}^*$$

$$\mathbf{w}(\delta(\varphi_{\beta\delta})) = v_{\beta}v_{\delta}v_{\delta}^*v_{\beta}^*$$

qui après saturation donnent :

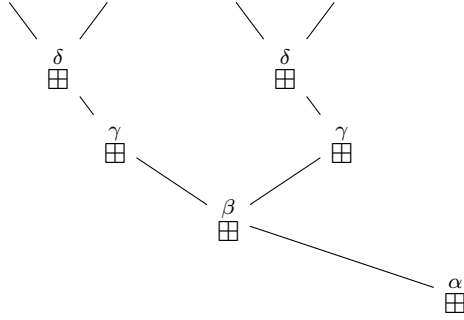
$$\overline{\mathbf{w}}(\delta(\varphi_{\alpha\gamma})) = u_{\alpha}u_{\gamma}e_{\alpha}e_{\gamma}u_{\gamma}^*u_{\alpha}^*$$

$$\overline{\mathbf{w}}(\delta(\varphi_{\alpha\delta})) = v_{\alpha}u_{\delta}e_{\alpha}e_{\delta}u_{\delta}^*v_{\alpha}^*$$

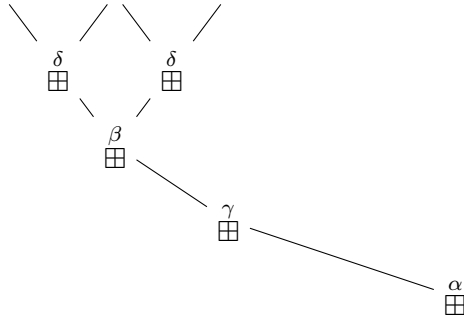
$$\overline{\mathbf{w}}(\delta(\varphi_{\beta\gamma})) = u_{\beta}v_{\gamma}e_{\beta}e_{\gamma}v_{\gamma}^*u_{\beta}^*$$

$$\overline{\mathbf{w}}(\delta(\varphi_{\beta\delta})) = v_{\beta}v_{\delta}e_{\beta}e_{\delta}v_{\delta}^*v_{\beta}^*$$

Remarquons que les deux réseaux simples finaux sont les seuls qui contiennent deux chemins après réduction. Si on considère le chemin $\delta^*(\varphi_{\alpha\gamma})$ il est composé de chemins dans des réseaux simples formant l'arbre de somme partielle :



qui est égal à l'arbre suivant dans notre quotient :



On remarque donc que tout chemin présent dans chacun des réseaux simples a la forme

$$\alpha^r \gamma^r (\beta^r \delta^r \varphi \beta \delta + \beta^r \delta^r \varphi \delta \beta + \beta^r \delta^r \varphi \beta \delta + \beta^r \delta^r \varphi \delta \beta) \gamma \alpha$$

et a donc le poids $u_{\alpha}u_{\gamma}\mathbf{w}(\varphi)u_{\gamma}^*u_{\alpha}^*$. Comme, de plus, α et γ ne sont plus présent dans les réseaux simples, on obtient nécessairement un poids de la forme $u_{\alpha}e_{\alpha}u_{\alpha}^*u_{\gamma}e_{\gamma}u_{\gamma}^*\mathbf{w}(\varphi)$ après normalisation.

Dans cet exemple on remarque donc que la tranche d'un chemin correspond à une branche de l'arbre de la forme normale telle que tout réduit du chemin passe par cette branche. De plus, si un nom apparaît dans l'arbre mais pas dans la tranche d'un chemin, celui-ci est nécessairement présent de part et d'autre des nœuds associés à ce nom.

11.4 Une réalisation de $\partial L^* : \partial \mathcal{L}$

11.4.1 Jetons

Soit Σ un ensemble fini, on note Σ^* l'ensemble des mots sur Σ et ϵ le mot vide. L'opération de concaténation sur les mots est noté $u \bullet v$.

En se restreignant sur les opérations que l'on effectue sur les mots, on manipule en fait des piles. Le mot $s_1 \bullet \dots \bullet s_n$, où $s_1, \dots, s_n \in \Sigma$ correspond à la pile contenant $s_1, \dots, s_n$, l'ajout d'une lettre à gauche correspond à l'empilement et l'effacement de la lettre la plus à gauche correspond au dépilement.

On note $F(A, B^*)$ l'ensemble des fonctions de A dans B^* telles que l'ensemble des x tels que $f(x) \in \epsilon$ soit fini. Soient $f \in F(A, B^*)$, $a \in A$ et $b \in B^*$, on note $f[a \mapsto b]$ l'élément de $\mathcal{IP}_{fin}(A, B)$ défini par

$$f[a \mapsto b](a') = \begin{cases} b & \text{si } a = a' \\ f(a') & \text{sinon} \end{cases}$$

Définition 11.28 *On appelle jeton un élément de*

$$\mathfrak{T} = \{0, 1\}^* \times (\{0, 1\}^*)^* \times (\{0, 1\}^*)^* \times F(\mathcal{N}, \{0, 1\}^*)$$

Pour $t = (M, E_?, E_!, f) \in \mathfrak{T}$ on appelle M la pile multiplicative, $E_?$ la pile ?-exponentielle, $E_!$ la pile !-exponentielle et f la fonction d'indexation

11.4.2 Opérations sur les jetons

On définit les opérations suivantes sur les jetons :

$$\begin{aligned} (M, E_?, E_!, f) &\xrightarrow{p} (0 \bullet M, E_?, E_!, f) \\ (M, E_?, E_!, f) &\xrightarrow{q} (1 \bullet M, E_?, E_!, f) \\ (M, \sigma \bullet E_?, E_!, f) &\xrightarrow{r_?} (M, (0 \bullet \sigma) \bullet E_?, E_!, f) \\ (M, \sigma \bullet E_?, E_!, f) &\xrightarrow{s_?} (M, (1 \bullet \sigma) \bullet E_?, E_!, f) \\ (M, E_?, E_!, f) &\xrightarrow{\partial_{\alpha, ?}} (M, \epsilon \bullet E_?, \sigma \bullet E_!, f) \text{ où } f(\alpha) = \sigma \\ (M, E_?, E_!, f) &\xrightarrow{u_{\alpha}} (M, E_?, E_!, f[\alpha \mapsto 0 \bullet \sigma]) \text{ où } f(\alpha) = \sigma \\ (M, E_?, E_!, f) &\xrightarrow{v_{\alpha}} (M, E_?, E_!, f[\alpha \mapsto 1 \bullet \sigma]) \text{ où } f(\alpha) = \sigma \\ (M, E_?, E_!, f) \text{ avec } f(\alpha) = \epsilon &\xrightarrow{e_{\alpha}} (M, E_?, E_!, f) \end{aligned}$$

Toutes ces opérations sont inversibles, et appartiennent donc à $\mathcal{IP}_{\mathfrak{T}}$.

Etant donné $\mathbf{g} \in \mathcal{IP}_{\mathfrak{T}}$ on peut la décomposer sous la forme $\mathbf{g} = (\mathbf{g}_m, \mathbf{g}_?, \mathbf{g}_!, \mathbf{g}_f)$ où chaque sous-opérations agit sur une composante du jeton. On peut alors définir l'opération duale $\Delta(\mathbf{g}) = (\mathbf{g}_m, \mathbf{g}_!, \mathbf{g}_?, \mathbf{g}_f)$. On pose alors $\mathbf{r}_! = \Delta(\mathbf{r}_?)$, $\mathbf{s}_! = \Delta(\mathbf{s}_?)$ et $\mathbf{d}_{\alpha, !} = \Delta(\mathbf{d}_{\alpha, ?})$.

11.4.3 Le modèle $\partial \mathcal{L}$

Définition 11.29 $\partial \mathcal{L}$ est le plus petit ensemble contenant les opérations définies précédemment et stable par composition et inversion. On note 1 l'identité et 0 l'opération nulle part définie.

Fait 11.30 $\mathfrak{d}\mathfrak{L}$ a une structure de monoïde inversif avec zéro. $0 \neq 1$ dans $\mathfrak{d}\mathfrak{L}$.

Preuve La structure est directement conséquence de la définition de $\mathfrak{d}\mathfrak{L}$. La différence entre 0 et 1 est due au fait que $\mathfrak{T} \neq \emptyset$. ✎

Théorème 11.31 $\mathfrak{d}\mathfrak{L} \models \partial L^*$

Corollaire 11.32 $0 \neq 1$ dans ∂L^*

Corollaire 11.33 $e_\alpha = e_\beta$ dans $\partial L^* \iff \alpha = \beta$

Preuve Pour les équations (11.3), (11.4) et (11.5), l'orthogonalité est déduite du fait que $0 \neq 1$ et, pour (11.3) et (11.5), l'orthogonalité pleine du fait que $\mathfrak{p}, \mathfrak{q}, \mathfrak{u}_\alpha$ et $\mathfrak{v}_\alpha$ ont $\mathfrak{T}$ pour domaine. L'équation (11.6) est triviale.

Les relations de commutation (11.7) sont issues du fait que $\mathfrak{r}_t$ et $\mathfrak{s}_t$ agissent uniquement sur la pile E_t ; les relations (11.8) de la focalisation de $\mathfrak{u}_\alpha$ et $\mathfrak{v}_\alpha$ sur la valeur de f en α .

Il est uniquement nécessaire de vérifier (11.9), l'équation (11.10) s'en déduisant par symétrie. On va en fait prouver l'inverse de l'équation, qui y est bien évidemment équivalente :

$$\mathfrak{r}_?^* \mathfrak{d}_{\alpha,!} = \mathfrak{u}_\alpha \mathfrak{d}_{\alpha,!} \mathfrak{u}_\alpha^*$$

Les opérations $\mathfrak{r}_?^* \mathfrak{d}_{\alpha,!}$ et $\mathfrak{u}_\alpha \mathfrak{d}_{\alpha,!} \mathfrak{u}_\alpha^*$ sont uniquement définies sur des jetons où $f(\alpha) = 0 \bullet \sigma$ et pour ceci on a :

$$\begin{aligned} & \mathfrak{r}_?^* \mathfrak{d}_{\alpha,!}(M, E_?, E_l, f) \\ = & \mathfrak{r}_?^*(M, (0 \bullet \sigma) \bullet E_?, \epsilon \bullet E_l, f) \\ = & (M, \sigma \bullet E_?, \epsilon \bullet E_l, f) \\ & \mathfrak{u}_\alpha \mathfrak{d}_{\alpha,!} \mathfrak{u}_\alpha^*(M, E_?, E_l, f) \\ = & \mathfrak{u}_\alpha \mathfrak{d}_{\alpha,!}(M, E_?, E_l, f[\alpha \mapsto \sigma]) \\ = & \mathfrak{u}_\alpha(M, \sigma \bullet E_?, \epsilon \bullet E_l, f[\alpha \mapsto \sigma]) \\ = & (M, \sigma \bullet E_?, \epsilon \bullet E_l, f) \end{aligned}$$

La dernière équation à vérifier est (11.11). On a

$$\begin{aligned} & \mathfrak{d}_{\alpha,!}^* \mathfrak{d}_{\beta,?}(M, E_?, E_l, f) \\ = & \mathfrak{d}_{\alpha,!}^*(M, \epsilon \bullet E_?, f(\beta) \bullet E_l, f) \\ = & (M, E_?, E_l, f) \end{aligned}$$

pour les jetons où $f(\alpha) = f(\beta) = \epsilon$; alors que l'opération $\mathfrak{e}_\alpha \mathfrak{e}_\beta$ est la projection sur les jetons satisfaisant cette condition. ✎

11.4.4 Une présentation alternative des jetons

La présence d'une fonction dans un jeton peut sembler étrange, cependant, la fonction d'indexation n'est rien d'autre que la représentation commode d'une infinité de piles indexées par $\mathcal{N}$. En fait une définition parfaitement équivalente d'un jeton aurait été celle d'un élément de

$$\mathfrak{T}' = \{0, 1\}^* \times (\{0, 1\}^*)^* \times (\{0, 1\}^*)^* \times_{\alpha \in \mathcal{N}} \{0, 1\}^*$$

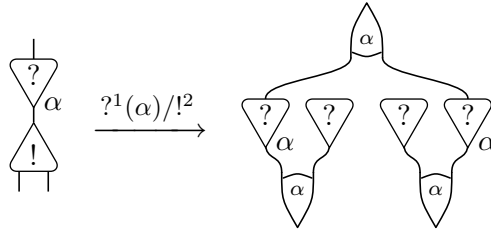
dont les piles sont presque toutes vides.

11.5 Graphes de superposition

Ce paragraphe présente un travail en cours. Nous avons cependant choisi de le présenter car il permet en l'état une meilleure compréhension des définitions précédentes et qu'il ouvre beaucoup de perspectives.

On va considérer ici de nouveaux réseaux que l'on appelle *graphes de superposition*, en référence aux graphes de partage. Du point de vue de leur statique, ils sont construits comme des réseaux d'interactions autour des cellules des réseaux d'interaction différentiels précédemment définies, en utilisant les variantes nommées des cellules de déréluction et co-déréluction, ainsi que d'une cellule binaire nommée $\triangleleft_\alpha$ appelée *cellule de superposition* et d'une cellule zéro-aire $\bullet$ appelée *cellule d'exception*.

Du point de vue de la dynamique, on remplace la réduction $?^1(\alpha)/!^2$ par

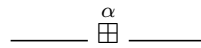


et on procède de même pour la règle duale. Pour tenir compte des nouvelles cellules introduites, on rajoute les règles données à la figure 11.3. La dernière règle présentée rend le système non normalisant. On peut en effet prouver la confluence, mais le fait que la règle soit involutive montre qu'elle est fortement divergente. On fera référence à cette règle sous le nom de *règle d'échange*.

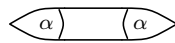
Un réseau différentiel se traduit naturellement en un graphe de superposition avec la construction suivante :

- on représente côte à côte tous les réseaux simples présents
- on choisit un représentant de l'arbre de somme
- pour chaque liste de port libre associés (rappelons que pour sommer des réseaux simples il faut mettre en relation leurs ports libres) on ajoute un arbre de cellules de superposition en reliant le port libre du réseau simple au port auxiliaire adapté

Par exemple le réseau simple



sera représenté par le graphe de superposition



Cette représentation n'est pas canonique car elle dépend d'un représentant de l'arbre de somme. Cependant, on peut toujours se ramener d'une représentation à une autre par réduction à l'aide la règle d'échange.

11.5.1 Exemples

Tout d'abord, reprenons l'exemple précédent du paragraphe 11.3.6. Après réduction on obtient le réseau normal suivant :

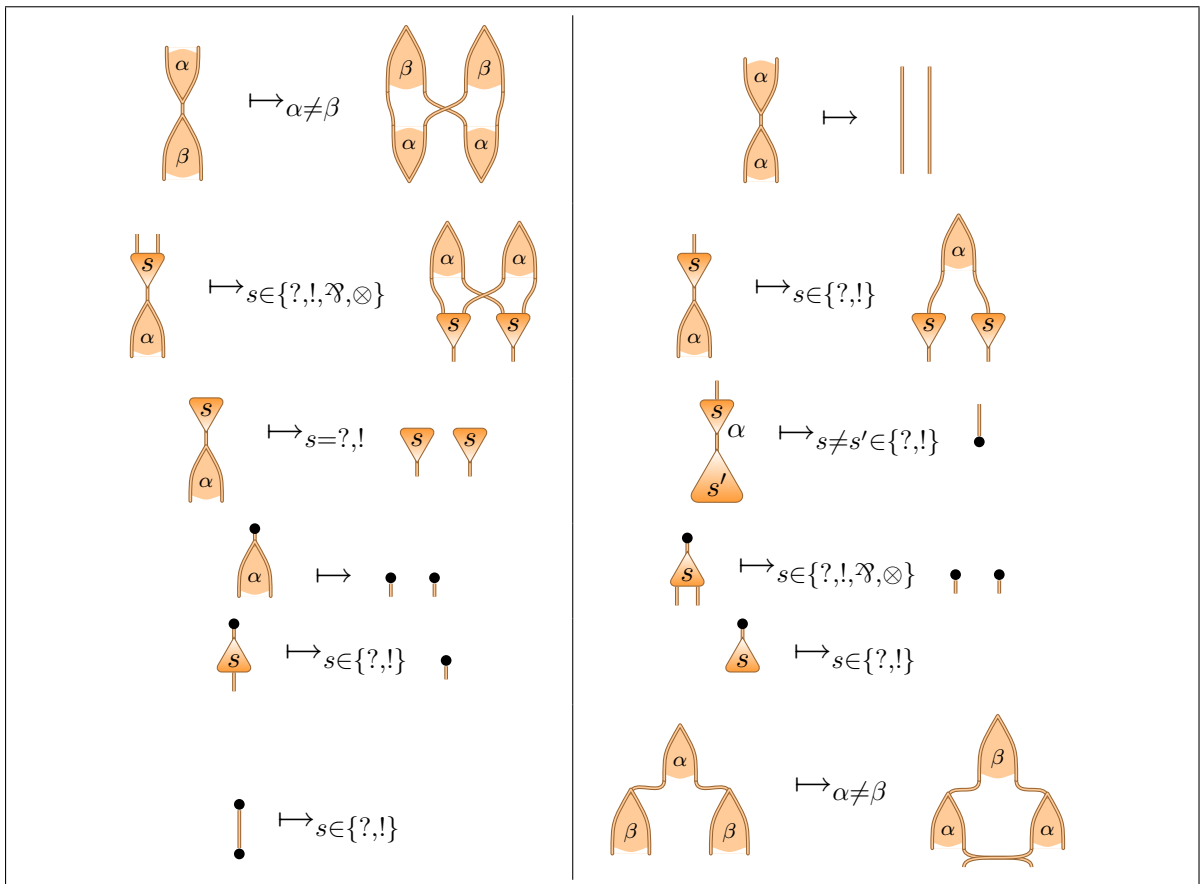
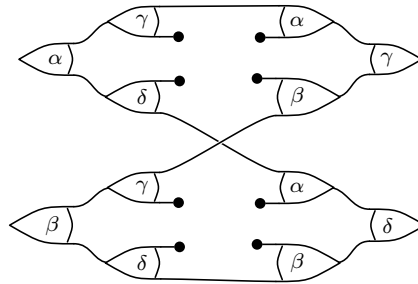
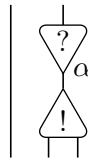


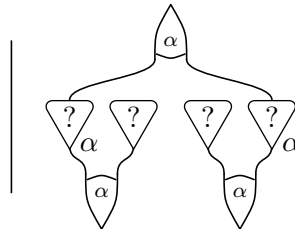
Figure 11.3: Les règles de réduction spécifiques aux graphes de superposition



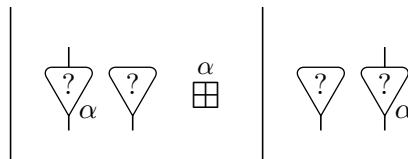
Si l'on compare ce graphe de superposition avec le réseau réduit présenté à la figure 11.2 on ne peut qu'être étonné de sa concision. En fait, ce graphe est une représentation fidèle de la GdI du réseau. Ce lien entre graphe de superposition et GdI est le même que pour les graphes de partage. On peut donc se demander si il existe une notion de réduction optimale qui leur correspond. A cet effet, réduisons le réseau suivant



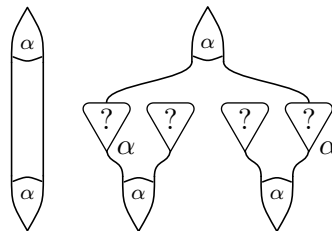
on obtient le graphe de superposition



alors que le réseau réduit est



dont le graphe de superposition est



On remarque ainsi que la réduction duplique bêtement le fil, alors qu'il n'a rien à voir avec le redex. En prenant le point de vue non-déterministe, c'est-à-dire celui où la somme exprime

un choix, il est clair que l'existence de ce fil est indépendante du choix α . Ainsi, dans tous les réduits on peut tenir compte de cette information et *superposer* les sous-réseaux égaux. Notons que pour que cette analyse soit fine il est nécessaire de restreindre certaines règles, celles liées à la propagation des cellules de superposition et ainsi à la duplication, pour obtenir une stratégie *paresseuse* visant à dupliquer le moins possible.

Perspectives

Dans chacune des parties nous avons développé des axes de recherche relativement indépendants. Nous présentons ici quelques perspectives pour chacun de ces axes.

Réseaux d'interaction multi-échelle

Suite à l'introduction des boîtes dans les réseaux d'interaction, il semble se dégager une structure plus générale de réseaux paramétriques. En effet, il est usuel de considérer des *encodages* de types de données par des sous-réseaux particuliers. Par exemple, les entiers unaires sont constitué d'un empilement de cellules unaire S (successeur) fermé par une cellule zéro-aire Z (zéro). Pour définir une opération sur ces entiers, on définit des règles de réduction qui préservent cette structure. Le problème dans cette approche est qu'elle repose entièrement sur une discipline de la part de celui qui définit les règles. Il serait beaucoup plus naturel de définir une cellule paramétré par un entier, qui pourrait être *déballée* progressivement en faisant apparaître la structure précédente. La définition de règles généralisées permettrait de caractériser celles préservant la structure particulière.

Une approche prometteuse consiste à considérer un réseau comme pouvant être vu à plusieurs échelles possibles. A l'échelle la plus atomique nous voyons un réseau standard, puis à mesure que l'on prend du recul on regroupe des paquets de cellules, faisant ainsi apparaître des sous-réseaux particuliers. La réduction pourrait aussi tenir compte de cette décomposition et permettre de tenir compte plus finement de la nature des réseaux.

Notons que cette approche, *les réseaux d'interaction multi-échelle*, a été partiellement implémentée et définie par nos soins. Il manque cependant une théorie fine permettant de faciliter la description et l'étude de ces réseaux, qui jusqu'à présent est assez laborieuse.

Concision de la GdI des λ -termes

Les différentes conjectures sur l'évolution de la GdI des λ -termes par la réduction, semblent converger vers une conjecture plus générale portant sur tout terme se réduisant vers un entier. Pour prouver cette conjecture, il est nécessaire de développer de nouvelles techniques de preuve. Nous pensons à ce propos qu'il faut dépasser le cadre de la théorie de la démonstration, et qu'une analyse algébrique plus poussée de la réduction est nécessaire.

GdI différentielle et concurrence

Ehrhard et Laurent ont défini dans l'article [EL07] une correspondance entre réseaux d'interaction différentiels et termes d'un π -calcul finitaire, préalable à une extension de l'isomor-

phisme de Curry-Howard à la concurrence. Mettant bout à bout cette correspondance et les techniques issues de la GdI des réseaux d'interaction différentiels, une notion pertinente de chemins dans des programmes concurrents apparaît naturellement. Cependant, l'utilisation très forte de l'associativité et de la commutativité de la somme, perdus par la GdI, pour établir cette correspondance oblige à développer des techniques algébriques pour compenser cela.

La structure algébrique des poids de chemins, dans les réseaux d'interaction différentiels, fait naturellement apparaître une sous-algèbre commutative liée aux choix non déterministes, permettant, par des techniques similaires aux algèbres de Boole, de caractériser l'union et l'intersection de différents choix. Ces techniques semblent donc permettre une spécification algébrique de comportements dynamiques complexes des réseaux, et par conséquent des programmes concurrents.

Par analogie avec la réduction optimale de Lévy, il semble se dégager une notion d'étiquettes pour les programmes concurrents, caractérisant ainsi leur évolution non-déterministe future. Une notion de stratégie de réduction optimale pourrait alors en découler, et permettre une meilleure compréhension des mécanismes de la concurrence.

Une construction standard permet de passer de la GdI du λ -calcul à la sémantique des jeux AJM. En effet, les ingrédients de la sémantique AJM sont les mêmes que ceux de la GdI, seul le cadre typé a priori change. Il est alors naturel d'essayer d'étendre cette construction pour donner une sémantique des jeux AJM dans le cadre des réseaux d'interaction différentiels. L'objectif étant, là encore, d'affiner la compréhension de la concurrence par des méthodes dynamiques issues de la théorie de la démonstration moderne.

Les analyses précédentes de la GdI des réseaux d'interaction différentiels et des programmes concurrents associés devraient permettre l'émergence d'une boîte à outils de techniques algébriques de spécification. Un objectif est d'en extraire un logiciel d'étude des programmes concurrents, fournissant ainsi, et de manière assez inédite, des techniques algébriques d'analyse. En effet, si de tels logiciels existent, ils reposent sur des techniques *ad hoc*, qui bien que très efficaces, ont l'inconvénient d'être dirigés par la syntaxe. L'approche issue de la GdI, et étant donc dirigée par la sémantique, pourrait les compléter.

Index

- affaiblissement, 51
- agrégation, 45
- algèbre d'un semigroupe inversif, 64
- algèbre stellaire, 64
- arbre exponentiel, 52

- β -réduction, 84
- bibliothèque, 32
 - close, 49
 - complète vis-à-vis du typage, 50
- boucle, 27
- boîte, 43
 - porte auxiliaire, 44
 - porte principale, 44

- cellule, 27
- charge, 82
- chemin, 67
 - dans des boîtes, 69
 - persistant, 75
 - régulier, 76
- co-affaiblissement, 105
- co-contraction, 105
- co-déréliction, 105
- concision, 91
 - de $\delta\bar{n}$, 91
- contexte, 30
- contraction, 51
- cycle, 19

- ∂L^* , 123
- duplication, 44
- déballage, 44
- découpage, 29
- déréliction, 51

- effacement, 44
- ELL, 52
- enfouissement, 45

- espace de Hilbert, 64
- Ex-composition, 22
- Ex-effondrement, 34
- Ex₀-composition, 21
- exécution, 19

- fil, 27
 - actif, 32
 - d'agrégation, 32
 - de partage, 32
 - libre, 32
- forme normale absolue, 33

- groupe, 58
- Géométrie de l'Interaction, 79
 - des réseaux différentiels, 115

- hauteur, 53

- idempotent, 57
- IN, 36
- injection partielle, 19
 - co-domaine, 19
 - disjointes, 19
 - domaine, 19
- interface, 30
- inverse, 57

- Krivine Abstract Machine, 85

- lambda-calcul
 - traduction dans SMELL, 86
- λ -calcul, 84
- λ -terme, 84

- MELL, 51
- MLL, 50
- modèle d'un monoïde inversif avec zéro, 60
- monoïde, 58
 - inversif, 58

-
- morphisme de réseaux d'interaction, 27
 - neutre, 57
 - orbite, 19
 - ordre, 19
 - paquet de constantes exponentielles, 82
 - permutation, 19
 - à orbites pointées, 19
 - étiquetée, 19
 - permutation partielle, 19
 - pondération, 75
 - fidèle, 76
 - fortement fidèle, 76
 - générique, 75
 - simple, 76
 - standard, 76
 - port, 26
 - auxiliaire, 27
 - de cellule, 27
 - libre, 27
 - principal, 27
 - portance, 53
 - promotion, 51
 - présentation par générateurs et relations, 59
 - pseudo-chemin, 69
 - recollement, 28
 - renommage, 27
 - règle d'interaction, 31
 - réduction de chemins, 70
 - avec boîtes, 72
 - réseau Axiome/Coupure, 33
 - réseau d'interaction, 26
 - support, 27
 - à boîtes, 43
 - réseau d'interaction différentiel, 105
 - réseau simple, 105
 - semigroupe, 57
 - ayant la propriété ab^* , 60
 - inversif, 58
 - normal, 59
 - positif, 59
 - simplifiable, 60
 - SMELL, 52
 - somme amalgamée, 37
 - sommes nommées de réseaux simples, 117
 - symboles, 26
 - théorème de Cayley, 62
 - théorème de Preston-Wagner, 63
 - typage, 49
 - pur, 89
 - w -permutation, 21
 - zéro, 57

Bibliographie

- [Abr97] S. Abramsky. Interaction, Combinators, and Complexity. *Lecture Notes, Siena, Italy*, 1997.
- [ACM90] Association for Computing Machinery. *Proceedings of the 17th Annual ACM Symposium on Principles of Programming Languages*, San Francisco, 1990. ACM Press.
- [ADLR94] Andrea Asperti, Vincent Danos, Cosimo Laneve, and Laurent Regnier. Paths in the lambda-calculus. In *Proceedings of the 9th Symposium on Logic in Computer Science*, Paris, 1994. IEEE Computer Society Press.
- [AG98] André Asperti and Stéfano Guerrini. *The Optimal Implementation of Functional Programming Languages*, volume 45. Cambridge University Press, 1998.
- [AGL92a] Martin Abadi, Georges Gonthier, and Jean-Jacques Lévy. The geometry of optimal lambda reduction. In *Proceedings of the 19th Annual ACM Symposium on Principles of Programming Languages*, pages 15–26. Association for Computing Machinery, ACM Press, 1992.
- [AGL92b] Martin Abadi, Georges Gonthier, and Jean-Jacques Lévy. Linear logic without boxes. In LICS'92 [LIC92].
- [AJ92] S. Abramsky and R. Jagadeesan. New foundations for the geometry of interaction. In LICS'92 [LIC92], pages 211–222.
- [AJM94] Samson Abramsky, Radha Jagadeesan, and Pasquale Malacaria. Full abstraction for PCF (extended abstract). In Masami Hagiya and John C. Mitchell, editors, *Theoretical Aspects of Computer Software. International Symposium TACS'94*, number 789 in Lecture Notes in Computer Science, pages 1–15, Sendai, Japan, April 1994. Springer-Verlag.
- [Ban95] R. Banach. The algebraic theory of interaction nets. *Department of Computer Science, University of Manchester, Technical Report MUCS-95-7-2*, 1995.
- [Dan90] Vincent Danos. *Une Application de la Logique Linéaire à l'Étude des Processus de Normalisation (principalement du λ -calcul)*. Thèse de doctorat, Université Paris 7, 1990.
- [dF08] Marc de Falco. The geometry of interaction of differential interaction nets. In *Proceedings of the 23th Symposium on Logic in Computer Science*, Pittsburgh, 2008. IEEE Computer Society Press.
- [dF09] Marc de Falco. An explicit framework for interaction nets. In Ralf Treinen, editor, *Rewriting Techniques and Applications (RTA '09)*, volume 5595 of *Lecture Notes in Computer Science*, pages 209–223. Springer, June 2009.

- [DR89] Vincent Danos and Laurent Regnier. The structure of multiplicatives. *Archive for Mathematical Logic*, 28 :181–203, 1989.
- [DR93] Vincent Danos and Laurent Regnier. Local and asynchronous beta-reduction. In *Proceedings of the 8th Symposium on Logic in Computer Science*. IEEE Computer Society Press, 1993.
- [DR95] Vincent Danos and Laurent Regnier. Proof-nets and the Hilbert space. In Girard et al. [GLR95].
- [Ehr02] T. Ehrhard. On Köthe sequence spaces and linear logic. *Mathematical Structures in Computer Science*, 12(05) :579–623, 2002.
- [Ehr05] T. Ehrhard. Finiteness spaces. *Mathematical Structures in Computer Science*, 15(04) :615–646, 2005.
- [EL07] Thomas Ehrhard and Olivier Laurent. Interpreting a finitary pi-calculus in differential interaction nets. In Luis Caires and Vasco T. Vasconcelos, editors, *Concurrency Theory (CONCUR '07)*, volume 4703 of *Lecture Notes in Computer Science*, pages 333–348. Springer, September 2007.
- [EPS73] H. Ehrig, M. Pfender, and HJ Schneider. Graph-grammars : an algebraic approach. In *IEEE Conference Record of 14th Annual Symposium on Switching and Automata Theory, 1973. SWAT'08*, pages 167–180, 1973.
- [ER03] Thomas Ehrhard and Laurent Regnier. The differential lambda-calculus. *Theoretical Computer Science*, 309(1-3) :1–41, 2003.
- [ER05] Thomas Ehrhard and Laurent Regnier. Differential interaction nets. In *Workshop on Logic, Language, Information and Computation (WoLLIC), invited paper*, volume 123 of *Electronic Notes in Theoretical Computer Science*. Elsevier, 2005.
- [FM99] M. Fernandez and I. Mackie. A calculus for interaction nets. *Lecture Notes in Computer Science*, 1702 :170–187, 1999.
- [Gir87] Jean-Yves Girard. Multiplicatives. In Lolli, editor, *Logic and Computer Science : New Trends and Applications*, pages 11–34, Torino, 1987. Università di Torino. Rendiconti del seminario matematico dell’università e politecnico di Torino, special issue 1987.
- [Gir89a] Jean-Yves Girard. Geometry of interaction I : an interpretation of system F . In Valentini Ferro, Bonotto and Zanardo, editors, *Proceedings of the Logic Colloquium 88*, pages 221–260, Padova, 1989. North-Holland.
- [Gir89b] J.Y. Girard. Towards a geometry of interaction. *Categories in Computer Science and Logic*, 92 :69–108, 1989.
- [Gir90] Jean-Yves Girard. Geometry of interaction II : Deadlock free algorithms. In Martin-Löf and Mints, editors, *Proceedings of COLOG'88*, volume 417 of *Lecture Notes in Computer Science*, pages 76–93. Springer-Verlag, Heidelberg, 1990.
- [Gir95] Jean-Yves Girard. Geometry of interaction III : accomodating the additives. In Girard et al. [GLR95].
- [Gir06] Jean-Yves Girard. Geometry of interaction IV : the feedback equation. In Stoltenberg-Hansen and Vaananen, editors, *Logic Colloquium 2003*. Association for Symbolic Logic, 2006.

-
- [Gir08] Jean-Yves Girard. Geometry of Interaction V : logic in the hyperfinite factor. *Theoretical Computer Science*, 2008.
 - [GLR95] Jean-Yves Girard, Yves Lafont, and Laurent Regnier, editors. *Advances in Linear Logic*, volume 222 of *London Mathematical Society Lecture Note Series*. Cambridge University Press, 1995.
 - [HO00] Martin Hyland and Luke Ong. On full abstraction for PCF. *Information and Computation*, 163 :285–408, 2000.
 - [How95] John M. Howie. *Fundamental of Semigroups Theory*. Oxford Press, 1995.
 - [HS06] E. Haghverdi and P. Scott. A categorical model for the geometry of interaction. *Theoretical Computer Science*, 350(2-3) :252–274, 2006.
 - [Hug05] Dominic J.D. Hughes. Simple multiplicative proof nets with units. Archived as math.LO/0507003 at arXiv.org. Submitted., March 2005.
 - [JSV96] A. Joyal, R. Street, and D. Verity. Traced monoidal categories, Math. In *Proc. Comb. Phil. Soc*, volume 119, pages 447–468, 1996.
 - [Laf90] Yves Lafont. Interaction nets. In ACM’90 [ACM90], pages 95–108.
 - [Laf03] Y. Lafont. Towards an Algebraic Theory of Boolean Circuits. *Journal of Pure and Applied Algebra*, 184(2-3) :257–310, 2003.
 - [Lam90] J. Lamping. An algorithm for optimal lambda calculus reductions. In ACM’90 [ACM90], pages 16–30.
 - [Lau01] Olivier Laurent. A token machine for full geometry of interaction (extended abstract). In Samson Abramsky, editor, *Typed Lambda Calculi and Applications ’01*, volume 2044 of *Lecture Notes in Computer Science*, pages 283–297. Springer, May 2001.
 - [LIC92] *Proceedings of the 7th Symposium on Logic in Computer Science*, Santa Cruz, 1992. IEEE Computer Society Press.
 - [Lip03] Sylvain Lippi. Encoding left reduction in the λ -calculus with interaction nets. *Mathematical Structures in Computer Science*, 12(06) :797–822, 2003.
 - [Mac95] Ian Mackie. The geometry of implementation. In *Proceedings of the 22th Annual ACM Symposium on Principles of Programming Languages*. Association for Computing Machinery, ACM Press, 1995.
 - [Mac98] I. Mackie. YALE : yet another lambda evaluator based on interaction nets. In *Proceedings of the third ACM SIGPLAN international conference on Functional programming*, pages 117–128. ACM New York, NY, USA, 1998.
 - [MP98] Ian Mackie and Jorge Sousa Pinto. Compiling the Lambda Calculus into Interaction Combinators. In *Logical Abstract Machines workshop*, 1998.
 - [Pin00] J. S. Pinto. Sequential and concurrent abstract machines for interaction nets. In *FOSSACS ’00 : Proceedings of the Third International Conference on Foundations of Software Science and Computation Structures*, pages 267–282. Springer-Verlag London, UK, 2000.
 - [Reg92] Laurent Regnier. *Lambda-Calcul et Réseaux*. Thèse de doctorat, Université Paris 7, 1992.
 - [Ser77] J.P. Serre. *Arbres, amalgames, SL_2* . Société mathématique de France, 1977.
 - [Vau07] Lionel Vaux. *Lambda-calcul différentiel et logique classique*. Thèse de doctorat, Université de la Méditerranée, 2007.