



Linear Logic and Sub-polynomial Classes of Complexity

Clément Aubert

► To cite this version:

Clément Aubert. Linear Logic and Sub-polynomial Classes of Complexity. Computational Complexity [cs.CC]. Université Paris-Nord - Paris XIII, 2013. English. NNT: . tel-00957653v2

HAL Id: tel-00957653
<https://theses.hal.science/tel-00957653v2>

Submitted on 9 Nov 2014

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.



Distributed under a Creative Commons Attribution - NonCommercial - ShareAlike 4.0 International License

Université Paris 13
Sorbonne Paris Cité
Laboratoire d'Informatique de Paris Nord

LOGIQUE LINÉAIRE ET CLASSES DE COMPLEXITÉ SOUS-POLYNOMIALES



LINEAR LOGIC AND SUB-POLYNOMIAL CLASSES OF COMPLEXITY

Thèse pour l'obtention du diplôme de
Docteur de l'Université de Paris 13,
Sorbonne Paris-Cité,
en informatique
présentée par Clément AUBERT
<aubert@lipn.fr>



Mémoire soutenu le mardi 26 novembre 2013, devant la commission d'examen composée de :

M.	Patrick BAILLOT	C.N.R.S., E.N.S. Lyon	<i>Rapporteur</i>
M.	Arnaud DURAND	Université Denis Diderot - Paris 7	<i>Président</i>
M.	Ugo DAL LAGO	I.N.R.I.A., Università degli Studi di Bologna	<i>Rapporteur</i>
Mme.	Claudia FAGGIAN	C.N.R.S., Université Paris Diderot - Paris 7	
M.	Stefano GUERRINI	Institut Galilée - Université Paris 13	<i>Directeur</i>
M.	Jean-Yves MARION	Lo.R.I.A., Université de Lorraine	
M.	Paul-André MELLIÈS	C.N.R.S., Université Paris Diderot - Paris 7	
M.	Virgile MOGBIL	Institut Galilée - Université Paris 13	<i>Co-encadrant</i>



* Créteil, le 7 novembre 2014 *

Errata

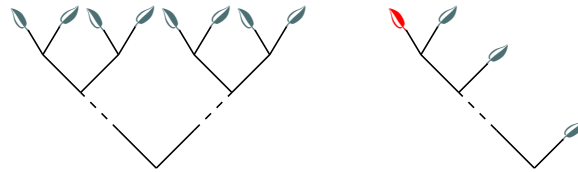
SOME minor editions have been made since the defence of this thesis (improvements in the bibliography, fixing a few margins and some typos): they won't be listed here, as those errors were not altering the meaning of this work.¹

However, while writing an abstract for [TERMGRAPH 2014](#), I realized that the simulation of Proof Circuits by an Alternating Turing Machine, in the [Section 2.4](#) of [Chapter 2](#) was not correct.

More precisely, $\mathbf{bPCC}^i$ is not an object to be considered (as it is trivially equal to $\mathbf{PCC}^i$) and [Theorem 2.4.3](#), p. 51, that states that “For all $i > 0$, $\mathbf{bPCC}^i \subseteq \mathbf{STA}(\log, *, \log^i)$ ”, is false. The ATM constructed do normalize the proof circuit given in input, but not within the given bounds. This flaw comes from a mismatch between the *logical depth* of a proof net ([Definition 2.1.10](#)) and the *height of a piece* ([Definition 2.2.6](#)), which is close to the notion of *depth of a Boolean circuit* ([Definition 2.1.2](#)).

First, remark that the class $\mathbf{bPCC}^i$ ([Definition 2.2.8](#)) is ill-defined: recall that ([Lemma 2.3.1](#)) in the case of $\mathcal{D}_{\text{isj}}^k$ and $\mathcal{C}_{\text{onj}}^k$ pieces, one entry is at the same logical depth than the output, and all the other entries are at depth 3 plus the depth of the output.

In the process of decomposing a single n -ary operation in $n - 1$ binary operations, the “depth-efficient way” is not the “logical depth-efficient way”. Let us consider a tree with n leaves (♣) and two ways to organize it:



Let us take the *distance* to be the number of edges between a leaf and the root of the tree. Then, on the left tree, every leaf is at the same distance $\log(n)$, and on the right tree, the greatest distance is n .

Everything differs if one consider that every binary branching is a $\mathcal{D}_{\text{isj}}^2$ or $\mathcal{C}_{\text{onj}}^2$ piece and consider the logical depth rather than the distance. In that case, on the left-hand tree, the i th leaf (starting from the left) is at logical depth 3 times the number of 1 in the binary encoding of i .² On the right-hand tree, every leaf is at logical depth 3, except for the red leaf (the left-most one), which is at logical depth 0.

So the logical depth is in fact independent of the fan-in of the pieces, and it has the same “logical cost” in terms of depth to compute $\mathcal{D}_{\text{isj}}^i$ or $\mathcal{C}_{\text{onj}}^i$ for any $i \geq 2$. A proof circuit of $\mathbf{PCC}^i$ is a proof circuit of $\mathbf{bPCC}^i$: every piece of $\mathcal{P}_u$ can be obtained from pieces of $\mathcal{P}_b$ *without increasing the logical depth*. So we have that for all $i \in \mathbb{N}$,

$$\mathbf{PCC}^i = \mathbf{bPCC}^i.$$

Secondly, the proof of [Theorem 2.4.3](#) was an attempt to adapt the function $\text{value}(n, g, p)$ [135, p. 65] that prove that an ATM can normalize a Boolean circuit with suitable bounds. The algorithm is correct,

¹Except maybe for a silly mistake in [Proposition 4.3.1](#), spotted by Marc Bagnol.

²Sic. This sequence is known as [A000120](#) and starts with 0, 1, 1, 2, 1, 2, 2, 3, 1, 2, 2, 3, 2, 3, 3, 4, ...

but not the bounds: it is stated that “Each call to f uses a constant number of alternations, and $\mathcal{O}(d(P_n))$ calls are made. As P_n is of depth $\log^i(n)$, $\mathcal{O}(\log^i(n))$ calls are made.”, and this argument is wrong.

Consider that $P_n \in \mathbf{PCC}^i$ (this is harmful taking the first remark into account), when the ATM calls f to evaluate the output of a $\mathcal{D}_{\text{isj}}^k$ or $\mathcal{C}_{\text{onj}}^k$ piece, it makes $k - 1$ calls to f at a lesser depth, and one call at the same depth. The only bound on the fan-in of the pieces, k , is the size of the proof circuit, that is to say a polynomial in the size of the input: this is a disaster.

tl;dr: $\mathbf{bPCC}^i$ should not be considered as a pertinent object of study and the correspondence between the proof circuits and the Alternating Turing Machines is partially wrong. This does not affect anyhow the rest of this work and this part has never been published nor submitted.

Créteil, 7 November 2014.



This thesis is dedicated to the free sharing of knowledge.

Cette thèse est dédiée au libre partage des savoirs.



Disclaimer: It is not possible for me to write these acknowledgements in a language other than French.

Remerciements. Il est d'usage de commencer par remercier son directeur de thèse, et je me plie à cet usage avec une certaine émotion : Stephano GUERRINI m'a accordé une grande confiance en acceptant de m'encadrer dans cette thèse en Informatique alors que mon cursus ne me pré-disposait pas tellement à cette aventure. Par ses conseils, ses intuitions et nos discussions, il a su m'ouvrir des pistes, en fermer d'autres, conseiller et pointer de fructueuses lectures.

Virgile MOGBIL a co-encadré cette thèse avec une attention et un soin qui m'ont touchés. Il a su m'accompagner dans tous les aspects de celles-ci, m'initier à la beauté de la recherche sans me cacher ses défauts, me guider dans le parcours académique et universitaire bien au-delà de ma thèse. Il su faire tout cela dans des conditions délicates, ce qui rend son dévouement encore plus précieux.

Patrick BAILLOT fût le rapporteur de cette thèse et comme son troisième directeur : il a su durant ces trois années me conseiller, me questionner, chercher à comprendre ce que je ne comprenais pas moi-même. C'est un honneur pour moi qu'il ait accepté de rapporter ce travail avant autant de justesse et de clairvoyance.

Je suis également ravi qu'Ugo DAL LAGO ait accepté de rapporter ce travail, et m'excuse pour l'avalanche de coquilles qui ont dû lui abîmer l'œil. Sa place centrale dans notre communauté donne une saveur toute particulière à ses remarques, conseils et critiques.

Je suis honoré qu'Arnaud DURAND — qui fût mon enseignant lorsque je me suis reconverti aux Mathématiques — ait accepté d'être le président de ce jury. La présence de Claudia FAGGIAN, Jean-Yves MARION et Paul-André MELLIÈS place cette thèse à l'intersction de communautés auxquelles je suis ravi d'appartenir. Je les remercie chaleureusement pour leur présence, leurs conseils, leurs encouragements.

Une large partie de ce travail n'aurait pas pu voir le jour sans la patience et la générosité de Thomas SEILLER, avec qui j'ai toujours eu plaisir à rechercher un accès dans la compréhension des travaux de Jean-Yves GIRARD.

§

Mon cursus m'a permis de rencontrer de nombreux enseignants qui ont tous, d'une façon ou d'une autre, rendu ce cheminement possible. J'ai une pensée particulière pour Louis ALLIX, qui m'a fait découvrir la logique et pour Jean-Baptiste JOINET, qui a accompagné, encouragé et rendu possible ce cursus. De nombreux autres enseignants, tous niveaux et matières confondus, m'ont fait aimer le savoir, la réflexion et le métier d'enseignant, qu'ils m'excusent de ne pouvoir tous les nommer ici, ils n'en restent pas moins présents dans mon esprit.

§

J'ai été accueilli au Laboratoire d'Informatique de Paris Nord par Paulin JACOBÉ DE NAUROS (qui encadrerait mon stage), Pierre BOUDES et Damiano MAZZA (qui ont partagé leur bureau avec moi), et je ne pouvais rêver de meilleur accueil. L'équipe Logique, Calcul et Raisonnement offre un cadre de travail amical et chaleureux dans lequel je me suis senti accompagné, écouté et bienvenu.

Devenir enseignant dans une matière que je connaissais mal (pour ne pas dire pas) aurait été une tâche insurmontable pour moi si elle n'avait eu lieu à l'IUT de Villeteuse, où j'ai trouvé une équipe dynamique et respectueuse de mes capacités. Laure PETRUCCI, Camille COTI, Jean-Michel BARRACHINA, Fayssal BENKHALDOUN et Emmanuel VIENNET m'ont tous énormément appris dans leur accompagnement patient et attentif.

Ce fût toujours un plaisir que de discuter avec mes collègues du L.I.P.N., ainsi qu’avec le personnel administratif de Paris 13. Brigitte GUÉVENEUX, tout particulièrement, sait marquer de par sa bonne humeur et son accueil chaleureux tout personne qui pénètre dans ces locaux, qu’elle en soit remerciée.

§

La science ne vit qu’incarnée, et je remercie Claude HOLL d’avoir tout au long des formations ALICE — aux intervenants exceptionnels — su placer la recherche et ses acteurs en perspective. Mes perspectives à moi ont souvent été questionnées, critiquées, par la presse libre (C.Q.F.D., Article XI, Fakir, La décroissance, Le Tigre, et bien d’autres), qui a su faire de moi un meilleur individu.

Mon engagement associatif s’est construit à côté, et parfois à contre-courant, de cette thèse, mais je sais aujourd’hui qu’il a été grandi par ce travail, qui a également profité de ces échanges. J’ai bien évidemment une pensée toute particulière pour microlab, mais je n’oublie par La Goutte d’ordinateur, l’écluse, et tous les acteurs rémois et parisiens de la scène *d.i.y.*

Le lien entre toutes ses dimensions a trouvé un écho dans les discussions avec les doctorants, tous domaines confondus, que j’ai eu la chance de croiser. Je tiens notamment à remercier et à saluer³ Aloïs BRUNEL, Sylvain CABANACQ, Andrei DORMAN, Antoine MADET, Alberto NAIBO, Mattia PETROLO, Marco SOLIERI, mais n’oublie aucun de ceux avec qui j’ai eu la chance de dialoguer.

Ma famille est passé par toutes sortes d’épreuves durant cette thèse, mais a toujours su m’écouter, me comprendre, m’encourager : me faire confiance et me soutenir.

Enfin, merci à Laura, pour notre amour.

Et merci à Olivier, pour — entre mille autres choses — ses sourires.



Et un merci spécial à l’équipe L.D.P. de l’Institut de Mathématiques de Luminy qui a su m’offrir un cadre serein où finaliser agréablement ce travail...

³Dans le « désordre alphabétique », comme l’écrivait Jean-Baptiste JOINET.

Tools. This thesis was typeset in the [URW Garamond No.8 fonts](#),⁴ with the distribution TeX Live 2014, using pdfTeX and the class Memoir.

The source code was created and edited with various [free softwares](#), among which [Git](#) who trustfully backedup and synced and [Texmaker](#).

I am greatly in debt to the [TeX - LaTeX Stack Exchange](#) community, and grateful to [Siarhei KHIREVICH](#)’s “Tips on Writing a Thesis in L^AT_EX”, that surely made me lose some time by pushing me forward on the rigorous typesetting.

The bibliography could not have been without [DBLP](#) and the drawings owe a lot to some of the examples gathered at [TeXample.net](#).



This work is under the [Creative Commons Attribution-NonCommercial-ShareAlike 4.0](#) or later licence: you are free to copy, distribute, share and transform this work; as long as you mention that this is the work of Clément AUBERT, you do not use this work for commercial use, you let any modification of this work be under the same or similar licence.

⁴Whose glyphs are digital renderings of fonts actually designed by Claude GARAMONT in the mid sixteenth century!

Contents

	Page
Errata	iii
Notations	xiii
Preface	xix
Introduction	I
I Historical and Technical Preliminaries	3
1.1 Computational Complexity Theory	4
1.2 Linear Logic	12
1.3 Previous Gateways between Linear Logic and Complexity	17
2 Proof circuits	21
2.1 Basic Definitions: Boolean Circuits and MLL_u	22
2.2 Boolean Proof Nets	28
2.3 Correspondence With Boolean Circuits	41
2.4 Correspondence With Alternating Turing Machines	45
3 Computing in the Hyperfinite Factor	57
3.1 Complexity and Geometry of Interaction	58
3.2 First Taste: Tallies as Matrices	60
3.3 Binary Integers	68
3.4 How to Compute	76
3.5 Nilpotency Can Be Decided With Logarithmic Space	83
4 Nuances in Pointers Machineries	93
4.1 A First Round of Pointers Machines	94
4.2 Finite Automata and L	98
4.3 Non-Deterministic Pointer Machines	106
4.4 Simulation of Pointer Machines by Operator Algebra	110
Conclusion and Perspectives	123
A Operator Algebras Preliminaries	125
A.1 Towards von Neumann Algebras	126
	ix

A.2	Some Elements on von Neumann Algebras	132
B	An Alternate Proof of $\text{co-NL} \subseteq \text{NPM}$	141
B.1	Solving a co-NL -Complete Problem	142
B.2	Pointer-Reduction	144
	Index	149
	Bibliography	151

List of Tables and Figures

	Page
CHAPTER 1	
1.1 The computation graph of an ATM	9
1.2 The rules of MELL	13
CHAPTER 2	
2.1 The components of a proof net: ax -link, $\otimes^n$ -link and $\wp^n$ -link	25
2.2 Unfolding a MLL_u proof net	27
2.3 Cut-elimination in parallel: t -, a - and m -reductions	27
2.4 An obstacle to cut-elimination in parallel	29
2.5 Adding or removing a crossing: how negation works	31
2.6 The conditional, the core of the computation	31
2.7 Branching variables on a Boolean proof net and evaluating it	32
2.8 Pieces of Proof Circuits	33
2.9 The improvement of “built-in” composition	36
2.10 Contractibility rules: relabelling and contracting a pseudo net	37
2.11 Composing two proof circuits	41
2.12 The simulation of a Boolean proof net by a Boolean circuit	44
2.13 The access to the input in an input normal form ATM	46
2.14 The translation from Alternating Turing Machines to proof circuits	48
2.15 How a value is guessed by an Alternating Turing Machine	49
2.16 A shorthand for the “guess and check” routine	50
2.17 A function to compute the value carried by an edge	52
CHAPTER 3	
3.1 The derivation of a unary integer, without contraction rule.	61
3.2 Unaries as proof net	62
3.3 Unaries as proof net, with contraction	64
3.4 Binary integers as derivations	70
3.5 The binary sequence $\star 001011 \star$ as a proof net and as a circular string	71
3.6 Some explanations about the matricial representation of a binary word	73
3.7 Representation of the main morphisms defined in the proof of Lemma 3.5.1	86
3.8 A representation of the computation of an operator	87

CHAPTER 4

4.1	How a KUM sorts integers	97
4.2	How to simulate Turing Machines with Finite Automata	105
4.3	The operators encoding the forward and backward move instructions	115
4.4	The operators encoding rejection	117

APPENDIX A

A.1	Dependency of the algebraic structures	127
-----	--	-----

APPENDIX B

B.1	The transition relation that decides <i>STConnComp</i>	143
-----	--	-----

Notations

The entries are divided between Complexity, Logic and Mathematics. Although some symbols are used in different ways, the context should always makes clear which denotation we have in mind.

We deal with classes of complexity separately, below this table of notations.

COMPLEXITY

ATM	Alternating Turing Machine.....	6
$\mathfrak{B}_b$	A set of (bounded) Boolean functions (a basis)	22
$\mathfrak{B}_u$	A set of (unbounded) Boolean functions (a basis)	22
$\mathbf{C}$	Family of Boolean circuit	23
$\mathbf{C}, \mathbf{C}_n$	Boolean circuits	22
$\mathbf{C}_M$	The set of pseudo-configurations of a $\text{PM}(p)$ M	108
$\text{DFA}(k)$	Set of deterministic Multi-Head Finite Automaton with k heads	99
$\text{DPM}(p)$	The set of deterministic pointer machines with p pointers.....	108
$G_{M(x)}$	Computation graph of an ATM M on input x	8
$G_{M,n}$	Computation graph of an ATM M on input of size n	46
$\widehat{(\cdot)}$	Translation from families of computation graphs to proof circuit families.....	47
$\mathcal{L}(M)$	Language accepted by a machine M	7
$\text{L}_{\text{DC}}(\mathbf{C})$	Direct Connection Language of a family of Boolean circuits $\mathbf{C}$	23
$\text{L}_{\text{DC}}(\mathbf{P})$	Direct Connection Language of a family of Boolean proof nets $\mathbf{P}$	32
$[\text{not } a]$	Any symbol of the alphabet except a	99
$\text{NFA}(k)$	Set of Non-Deterministic Multi-Head Finite Automaton with k heads.....	98
$\text{NPM}(p)$	The set of non-deterministic pointer machines with p pointers	108
$\mathcal{O}$	“Big O” asymptotic notation	5

$\text{PM}(p)$	The set of (deterministic and non-deterministic) pointer machines with p pointers . 108
$P_{+,1}$	Set of boolean observations of 1-norm ≤ 1 83
P_+	Set of boolean observations 83
$P_{\geq 0}$	Set of positive observations 83
$Q^\dagger$	Extended set of states 110
$\#H$	Distance between $\triangleright$ and the bit currently read by the head H 99
σ	Transition relation of a NDPM. 98
$ x $	The size of x 5
$\star$	First and last bit of a circular binary word. 69
$\triangleright$	Left end-marker on the input tape of a FA. 98
$\triangleleft$	Right end-marker on the input tape of a FA 98
$\vdash_{M(x)}$	The transition relation between configurations in an ATM M on input x 7
$\langle 1^y, \bar{g}, \bar{p}, \bar{b} \rangle$	The tuples used to describe a proof net. 32
$\{X\}$	Language decided by the set of observations X 82

LOGIC

$A[B/D]$	The substitution of every occurrence of B by D in A 24
$A[D]$	The substitution of every occurrence of α by D in A 24
α	Literal (or atom). 13
$\overrightarrow{A}, (\text{resp. } \overleftarrow{A})$	The ordered sequence of formulae $A_1, \dots, A_n$, (resp. $A_n, \dots, A_1$) 24
$\Rightarrow$	The parallel evaluation of a proof net 28
$\mathbf{B}$	The Boolean type of MLL_u 28
b_0, b_1	Proof nets representing 0 and 1 in MLL_u 28
$(\cdot)^\perp$	Negation, or duality 13
$\mathbf{P}$	Boolean proof net 30
$\mathbf{P}$	A Boolean proof net family 30
$d(A)$	Depth of a MLL_u formula A 25
$\rightarrow_{\text{ev.}}$	The “evolution” through normalisation of a Boolean proof net 30

$\vdash \Gamma$	A sequent, Γ being a multiset of formulae.....	13
$\Rightarrow$	Classical implication	12
$\multimap$	Linear implication	13
MLL_u	Multiplicative Linear Logic	24
$\bullet$	An axiom link, in a proof net	25
$!$	The “of course” modality	13
$\wp$	The multiplicative disjunction (“par”).....	13
$\mathcal{P}$	Piece of a proof circuit	33
$\mathcal{P}_b$	Set of pieces of bounded fan-in	35
$\mathcal{P}_u$	Set of pieces of unbounded fan-in	35
$\mathcal{P}, \mathcal{Q}$	Proof nets of MLL_u	25
π	A proof, in the sequent presentation	14
$\otimes$	The multiplicative conjunction (“tensor”).....	13
$?$	The “why not” modality	13

MATHEMATICS

α	Action of a group.....	137
$ \cdot $	Absolute value	132
$\mathcal{B}(\mathbb{H})$	Banach algebra of continuous mapping over $\mathbb{H}$	131
$\det(M)$	Determinant of a matrix M	134
d	Distance function	126
δ_{ij}	Kronecker delta	126
$\sim_{\mathfrak{M}}$	Murray and von Neumann equivalence over the von Neumann algebra $\mathfrak{M}$	133
$\mathcal{F}$	Field.....	127
$\lfloor \cdot \rfloor$	Floor function	102
$\mathcal{G}$	Group	127
$\mathbb{H}$	Hilbert space.....	130
$\ell^2(S)$	Square-summable functions over S	132

$M[i, j]$	Element in the i th row and j th column of M	134
$M_{n,m}(\mathcal{F})$	Matrix of dimensions $n \times m$ with coefficients in the field $\mathcal{F}$	134
$\mathfrak{M}$	von Neumann algebra.....	131
$\mathfrak{M} \rtimes \mathcal{G}$	Crossed product of a von Neumann algebra $\mathfrak{M}$ with a group $\mathcal{G}$	137
$\ \cdot\ _1^k$	The 1-norm.....	82
$\pi_{k,i}$	Projection of dimension $(k+1) \times (k+1)$	74
$\Pi(\mathfrak{M})$	Set of projections of a von Neumann algebra $\mathfrak{M}$	133
$\mathcal{P}_{\text{fin}}(S)$	Set of finite subsets of S	6
$\mathfrak{R}$	The II_1 hyperfinite factor.....	134
$(\cdot)^*$	Involution	130
$\mathcal{S}(S)$	Permutation group of S	128
$\mathcal{S}$	Group of all finite permutations over $\mathbb{N}$	128
$\text{sgn}(\sigma)$	Signature of a permutation σ	128
$\otimes$	Tensor product	133
$\mathcal{T}$	Topology (over a set).....	131
$\text{tr}(M)$	Trace of a matrix M	134
$\text{Tr}(M_n)$	Normalised trace of a $n \times n$ matrix M_n	135
$\mathcal{G}$	Underlying set of the algebraic structure $\mathcal{G}$	127
V	Vector space	128

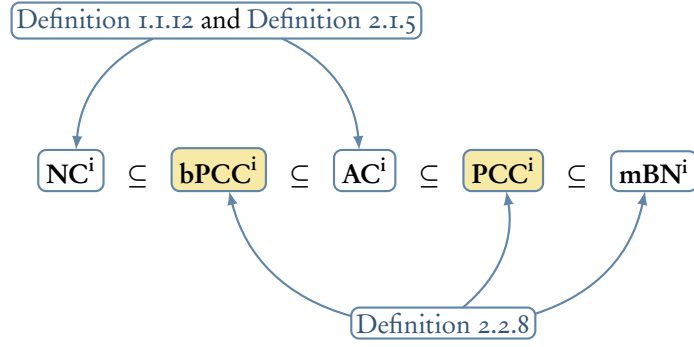
See referenced pages for formal definitions.

We are going to use several classes of complexity, some of them classical, some of them being introduced in this work (in a yellow rectangle in the figure below). The following pictures describe the situation at the end of our work. Some equalities and inclusions are classical, but we will in the course of this work give at least ideas of proofs for every of them.

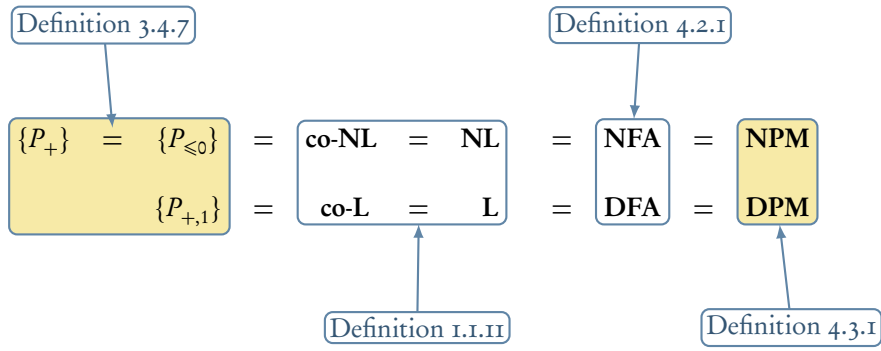
For all $i > 1$,⁵ we have at the end of [Chapter 2](#)⁶ the following situation:

⁵Even if those bounds are going to be sharpened, we prefer to give our results in all generality at this point.

⁶See errata.



Chapter 3 and Chapter 4 will define and prove properties about the following classes:



We are also going to introduce $STA(*,*,*)$ (Definition 1.1.7) as a convenient notation to express any of these classes.

Preface

THIS work is somehow two-sided, for it studies two models of computation —and the implicit characterization of complexity classes that comes with— that share the same tool and the same aim, but are not directly related. The tool —mainly proof theory, and more precisely Linear Logic— and aim —to contribute to the complexity theory— are presented in [Chapter 1](#). The [Section 1.3](#) presents some of the previous attempts to link those two fields.

The first model —presented in [Chapter 2](#)— is related to the simulation of efficient computation in parallel with Boolean proof nets. It introduces proof circuits as an innovative way of understanding the work led by TERUI. It sharpens the characterization of some complexity classes, and provides new evidence that this model is pertinent by linking it to the more classical Turing Machines. This chapter is —almost— self-contained and a gentle introduction to some important concepts.

The second model is presented in [Chapter 3](#), it develops a new world where logic is dynamically built from von Neumann algebras. We tried to ease the reading of this quite technical material by pushing slowly the constructions and taking time to explain some of the choices we made. We finish this first part by showing how complex our object is from a complexity point of view.

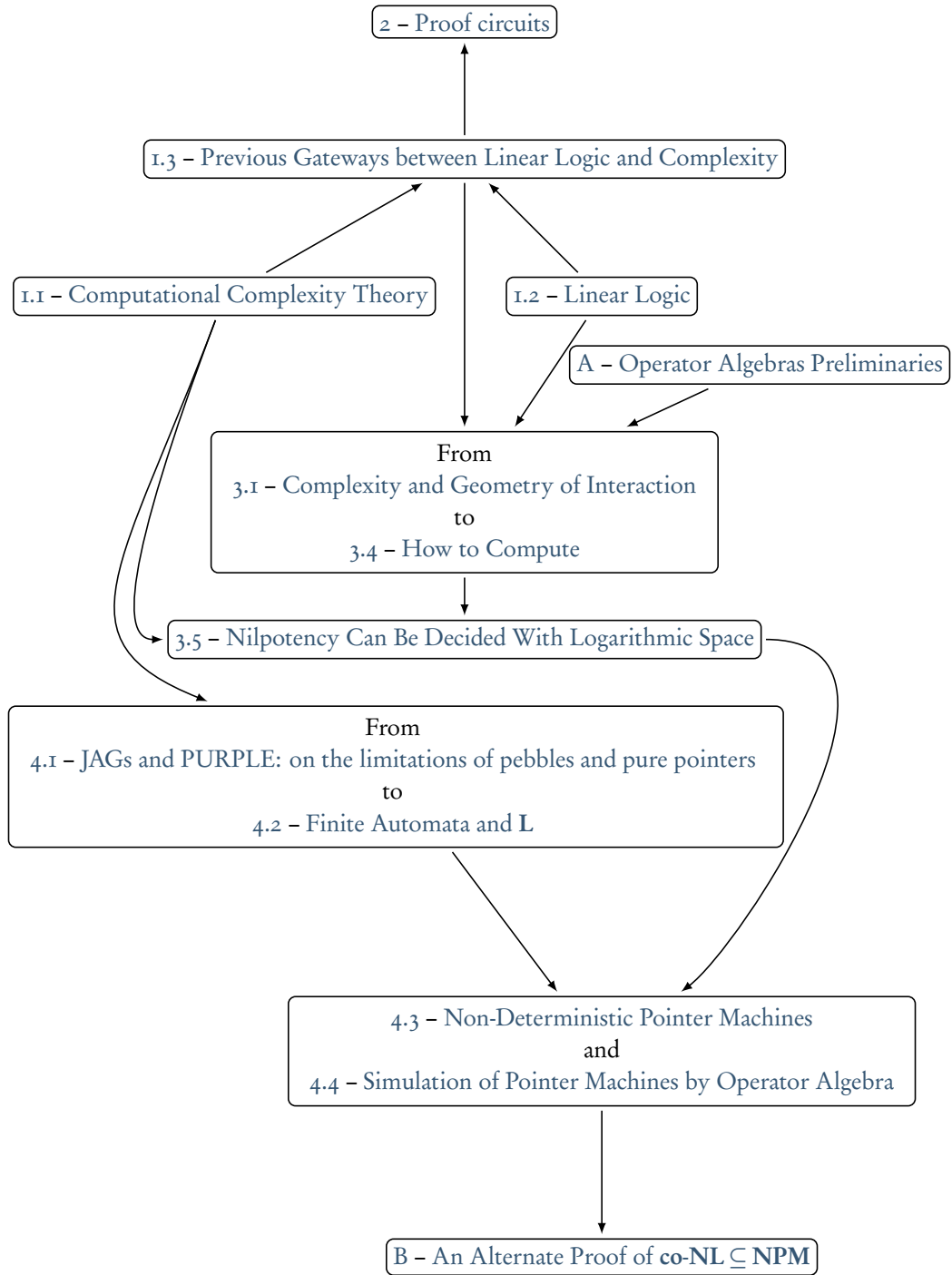
[Chapter 4](#) gently presents some “pointer machines” before defining a model that is proven adapted to be simulated by the model of [Chapter 3](#). It recollects some historical variation to help the reader to understand the framework we are working with.

We made the choice to put aside in [Appendix A](#) the introduction of the mathematical tools needed to understand the most advanced elements of [Chapter 3](#). This disposition could disturb the reader by forcing him/her to go back and forth between [Chapter 3](#) and this appendix. Yet, those tools should remain only tools, for they are not what is at stake here, and we tried to go as rapidly as possible to the point without getting too technical.

The [Appendix B](#) gives an alternate proof of a statement of [Chapter 4](#). It is somehow more algorithmic, for two problems are proved solvable by machines that are actually described. It is worth reading for it gets some insight on pointer computation, and by extension on log-space reduction.

The picture that follows presents some of the dependencies between the different parts of this work.





Introduction

LOGIC is so tightly bound to computer science that the sonship is somehow reversed: computers, through programming languages, now teach us how reasoning can be algorithmically represented. We learnt how to modelise the cost of solving a problem, answering a question mathematically presented. And yet Logic, thanks to the proof-as-program correspondence, still has a lot to say about this modelisation.

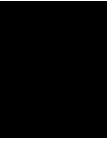
The first cost models that were developed on abstract machines discriminate between the “reasonable” models of computation and the one that are not. Of course, they could not foresee the inherent complexity of solving a precise problem: one had to actually run the program to observe the resources it needs to obtain a result. Some important problems, in a way equivalent to any other “as complex” problem, were defined and were taken as representative of the power of some complexity classes.

The Implicit Computational Complexity approach reversed this dependency, by defining frameworks where any algorithm that could be expressed was “pre-bounded”. Among the approaches developed stands Linear Logic, and among the classes “captured” stands $\mathbf{P}$, the class of “tractable” problems. This class drew a large part of the attention because of the challenge its separation from the other classes represents.

There is no reason to stick to that class, for Linear Logic has at its disposal numerous tools to simulate other ways of computing. Modalities —the salt of Linear Logic— were proven fruitful to handle the polynomial time complexity, but there are numerous other built-in mechanisms to cope with the complexity.

This work presents two innovative ways to represent non-deterministic and parallel computation with Linear Logic. They were not invented by the author, but by TERUI and by GIRARD, respectively in 2004 and in 2012. Those approaches have some elements in common, at least a part of the tools and the aims, and they are two natural extensions of a tradition of bridging the gaps between Linear Logic and complexity. Still, they are largely unexplored, excepted for two works led by MOGBIL and RAHLI, and except of SEILLER’s PhD Thesis.

This work may be hard to get through, for it mixes three fields of research: complexity theory, Linear Logic, and a mathematical framework known as Geometry of Interaction. We tried to ease the presentation by developing some examples and by taking some time to discuss the general shape of the construction as well as some details. May this work ease the understanding of those two exciting approaches.



Historical and Technical Preliminaries

THIS chapter will set up most of the elements needed to understand the rest of this work. A strict minimal knowledge on set theory and graph theory is required to grasp some definitions, but apart from that everything should be quite self-explanatory. However, we do not always take the time to introduce gently the notions and go straight to the point, but many references are provided.

This work takes place in the complexity area, a subject that mix theoretical computer science and mathematics, and tries to classify computational problems by their inherent complexity. This was historically done by showing off algorithms that run on abstract machines, where every operation is endowed with a cost, and by studying the overall cost of the computation. To be pertinent, this consumption of resources is expressed *relatively to the size of the input* by a function that does not have to be precise: we ignore the constant factors and the lower order terms to focus on a larger magnitude.

Later on came the *Implicit Computational Complexity*, that describe complexity classes without explicit reference to a machine model and to bounds. It massively uses mathematical logic to define programming language tools —like type-systems — that enforce resource bounds on the programs. Among the techniques developed, one could mention

- Recursion Theory, which studies the class characterized by restricting the primitive recursion schema.¹
- Model Theory, which characterizes complexity classes by the kind of logic needed to express the languages in them. It is also known as Descriptive complexity, and we refer to IMMERMANN [77] for a clear and self-contained presentation.
- Proof Theory, which uses the Curry-Howard correspondence to match up execution of a program and *normalisation* of a proof.

This latter way of characterizing complexity classes by implicit means is a major accomplishment of the cross-breeding of Logic and computer science. We will in the third section gives some examples of the previous works led onto that direction, but we will first present its two components: computational complexity theory, and a special kind of logic, Linear Logic.

¹The seminal work was made by LEIVANT, but one could with benefits look to a clear and proper re-statement of his results by DAL LAGO, MARTINI, and ZORZI [34].

1.1 Computational Complexity Theory

A short and yet efficient introduction to the theory of complexity can be found in the chapters “Basic notions in computational complexity” written by JIANG, LI, and RAVIKUMAR [82] and published in a classical handbook[6]. In the same book, the two chapters written by ALLENDER, LOUI, and REGAN [1, 2], “Complexity Classes” and “Reducibility and Completeness” synthesise all the results that will be used in this section. The handbook of PAPADIMITRIOU [111] is often refereed to as “the” handbook for a nice and complete introduction to this theory: its unique style allows to read it as a novel, and to grasp the tools one after the other, but for the very same reason it is hard to look for a precise information in it. The reader searching for a precise definition or a theorem should probably rather refer to the textbooks of GREENLAW, HOOVER, and RUZZO [65], and ARORA and BARAK [4].

Complexity and models of computation

This section will quickly try to present what a “reasonable” model of computation is, what is at stake when we take parallel models of computation, and gives hints toward the specificities of subpolynomial classes of complexity. It borrows some elements from VAN EMDE BOAS [132], whose preliminary chapter is clear and concise.

To study complexity in an “explicit” framework, one needs a machine model and a cost model. A machine model is a class of similar structures, described set-theoretically as tuples most of the time, endowed with a program, or finite control, and a definition of how *computation* is actually performed with a machine of that model. The cost model fixes the cost of the basic operations, allowing to measure the resources (most of the time understood as time and space) needed to process an input.

Two fundamentals of this theory is that something is “algorithmically computable” iff it is computable by a Turing Machine (Church–Turing thesis), and that taking the Turing Machine or another “reasonable” model of computation does not really matter, this is the invariance thesis:

Invariance Thesis. *“Reasonable” machines can simulate each other within a polynomially bounded overhead in time and a constant-factor overhead in space.*

Of course, what is a simulation needs to be define.

Definition 1.1.1 (Simulation). Given two machines models M and M' , we say that M' simulates M with time (resp. space) overhead $f(n)$ if there exists two functions c and s such that, for every machine M_i in M :

- for all input x of M , $c(x)$ is the translation of x in the setting of M' ;
- $M_i(x)$ halts iff $M'_{s(i)}(c(x))$ halts, and $M_i(x)$ accepts iff $M'_{s(i)}(c(x))$ accepts²;
- if $t(|x|)$ is the time (space) bound needed by M_i for processing x , then the time (resp. space) required by $M'_{s(i)}$ for processing $c(x)$ is bounded by $f(t(|x|))$;

If the time overhead function f is linear, the simulation is said to be *real-time*.

²Where $M_i(x)$ denotes the processing of x by M_i .

Something is missing in this “definition”: what are the resources needed to compute s and c ? It would be meaningless to provide more computational power to compute $M'_{s(i)}$ than what $M'_{s(i)}$ can actually provide: we could as well have computed the output directly! In other words, if the machine used to process c and s is strictly more powerful than $M'_{s(i)}$, we cannot extract any knowledge about M' . In the following, we need to be careful about the cost of those simulations.

We will tend to keep the word simulation for the case where c and s are computed by a machine of M' itself: in other word, provided a description of M_i , the machine M'_i will by itself “acts as M_i ”, without external intervention. If a third machine model M'' , given the description of M_i and x , outputs the description of $M'_{s(i)}$ and $c(x)$, we will speak of a *translation*.

Notice that f takes as input $|x|$ the size of the input, for we measure the resources needed by the computation *relatively to the size of the input*. This function should be *proper*,³ i.e. it should respect some basic properties, but we won't take that parameter into account, for we will use only the proper functions below:

Polynomial i.e. there exists $c, d \in \mathbb{N}^*$ such that $f(|x|)$ is bounded by $c \times (|x|^d)$.

Constant i.e. there exists $c \in \mathbb{N}^*$ such that $f(|x|)$ is bounded by $c \times |x|$.

(Binary) Logarithm i.e. there exists $c \in \mathbb{N}^*$ such that $f(|x|)$ is bounded by $c \times (\log(|x|))$.

Poly-logarithmic i.e. there exists $c, i \in \mathbb{N}^*$ such that $f(|x|)$ is bounded by $c \times (\log^i(|x|))$.

We will use the “big O” asymptotic notation, i.e. we write that $g = \mathcal{O}(f)$ iff $\exists k \in \mathbb{N}^*$ such that $\forall x, g(x) \leq k \times f(x)$. So if f is a polynomial (resp. the constant function, the logarithm, a poly-logarithm), we will sometimes write $n^{\mathcal{O}(1)}$ (reps. $\mathcal{O}(1)$, $\mathcal{O}(\log)$, $\mathcal{O}(\log^i)$) or even more simply poly (resp. 1, log, $\log^i$), without any reference to the size of the input.

This notation is a bit counter-intuitive, because despite the equal sign, it is not symmetric: $\mathcal{O}(x) = \mathcal{O}(x^2)$ is true but $\mathcal{O}(x^2) = \mathcal{O}(x)$ is not. As SIPSER [124, p. 277] writes it, “[t]he big-O interacts with logarithms in a particular way”: the changing of the base of the logarithm affects the complexity only by a constant factor, so we may simply ignore it and we will consider that the logarithmic function is always in base 2. We will use in the course of this work some classical identities, but maybe it is worth remarking that as $n = 2^{\log(n)}$, for all $d \in \mathbb{N}$, $n^d = 2^{d \log(n)}$ and so $2^{(\log(n))^{\mathcal{O}(1)}} = n^{\mathcal{O}(1)}$. We can also remark that as $\log(n^d) = d \log(n)$, $\log(n^{\mathcal{O}(1)}) = \mathcal{O}(\log(n))$. This will be more or less the only two equalities we are going to use.

We will focus here on the subpolynomial world, that is to say the machines that use less resources than polynomial time. This world is less studied than the polynomial word, as the “reasonable” computation is quite limited in it, but a fruitful approach often used is to look toward parallel computation. This world also have a kind of invariance thesis to define what is a reasonable model:

Parallel Computation Thesis. *Whatever can be solved in polynomially bounded space on a reasonable sequential machine model can be solved in polynomially bounded time on a reasonable parallel machine, and vice versa.*

This thesis was first formulated by GOLDSCHLAGER [64], and by CHANDRA, KOZEN, and STOCKMEYER [24] quite simultaneously. The first one stated that “Time-bounded parallel machines are polynomially related to space-bounded computers” by trying to “capture mathematically this intuitive

³For a precise definition, the reader should for instance refer to PAPADIMITRIOU [111, p. 140, definition 7.1.].

notion of parallel computation, much in the same spirit that ‘Church’s thesis’ captures the notion of effective computation.” CHANDRA, KOZEN, and STOCKMEYER, on their way, defined *Alternating Turing Machine* and proved that “an alternating Turing machine is, to within a polynomial, among the most powerful types of parallel machines”.

Alternating Turing Machines are also the most general version of the most well-known model, Turing Machines, so it is quite natural to take them as the reference, as the basic model. We should make one last remark concerning this parallel world: time is sometimes refereed to as the *depth*, i.e. the maximal number of intermediate steps needed ; and space is taken to be the *size*, i.e. the number of processors needed to perform the computation. This will be the case with the Boolean circuits and the proof circuits, defined in Chapter 2.

Alternating Turing Machine

This model of computation is really handy, for it will allow us to define at the same time parallel and sequential computing, and it has a characterization of every complexity class we will use in the following. The following definition borrows some elements of the clear presentation of VOLLMER [135, p. 53].

Definition 1.1.2 (ATM, CHANDRA, KOZEN, and STOCKMEYER [24]). An *Alternating Turing Machine* (ATM) M with $k \in \mathbb{N}^*$ tapes is a 5-tuple

$$M = (Q, \Sigma, \delta, q_0, g)$$

with

- Q a finite set of *states*,
- Σ an alphabet,⁴
- $\delta : (Q \times \Sigma^{k+1} \times \mathbb{N}^{k+1}) \rightarrow \mathcal{P}_{\text{fin}}(Q \times \Sigma^k \times \{-1, 0, +1\}^{k+1})$ the *transition relation*,
- $q_0 \in Q$ the *initial state*,
- and $g : Q \rightarrow \{\forall, \exists, 0, 1\}$ the *state type function*.

This definition does not explain how one can actually perform a computation with ATMs, and we will quickly expose those —very classical— mechanism right after we introduce some nomenclature.

The input is written on the *input tape*, and the head that scans it is read-only. The other k heads are read-write, and each one of them work on a dedicated tape. A *configuration* α of M is an element of $Q \times (\Sigma^*)^k \times \mathbb{N}^{k+1}$ that entirely describe the current situation. It reflects the current state, the content of the k working tapes and the position of the $k + 1$ heads, as a “snapshot” from where it is possible to resume the computation, provided that we know the input and the transition relation.

The initial configuration α_0 of M is when the entry is written on the input tape, the k working tapes are blank,⁵ the $k + 1$ heads are on the first cell of their tape, and the state is q_0 .

A state $q \in Q$ —and by extension a configuration whose state is q — is said to be *universal* (resp. *existential*, *rejecting*, *accepting*) if $g(q) = \forall$ (resp. $\exists$, 0, 1). A state $q \in Q$ —and a configuration whose state is q — is *final* if $\forall a_0, \dots, a_k \in \Sigma, \delta(q, a_0, \dots, a_k) = \emptyset$. This condition is equivalent to $g(q) \in \{0, 1\}$.

⁴One should precise that every tape is divided into *cells* and that every *letter* of the alphabet need exactly one cell to be written.

⁵Usually a special “blank” symbol $\flat$ is written on every cell, but it works as well to take any symbol of the alphabet.

Suppose that M is working on input $x \equiv x_1 \dots x_n$ (we denote this situation with $M(x)$) and that at a certain point, M is in configuration

$$\alpha = (\mathbf{q}, w_1, \dots, w_k, n_0, \dots, n_k)$$

where $w_i \equiv w_{i,1} \dots w_{i,l_i}$ for $1 \leq i \leq k$. Suppose that $\delta(\mathbf{q}, x_{n_0}, w_{1,n_1}, \dots, w_{k,n_k})$ contains (possibly among others) the tuple $(\mathbf{q}', a_1, \dots, a_k, X_0, \dots, X_k)$, where $a_1, \dots, a_k \in \Sigma$ and $X_0, \dots, X_k \in \{-1, 0, +1\}$. Then

$$\beta = (\mathbf{q}', w'_1, \dots, w'_k, n'_0, \dots, n'_k)$$

is a *successor configuration* of α , where for $1 \leq i \leq k$, w'_i is obtained from w_i as follow: replace the n_i th letter by a_i , and to cover the case where the head moves to the left or right out of the word by prefixing or suffixing the blank symbol $\flat$; and for $0 \leq i \leq k$, let $n'_i = n_i + X_i$, if $n_i + X_i \geq 0$, otherwise we leave $n'_i = n_i$.

That β is such a successor configuration of α (and respectively that α is a *predecessor* of β) is denoted by $\alpha \vdash_{M(x)} \beta$, or if M and x are clear from the context by $\alpha \vdash \beta$. As usual, $\vdash^*$ is the reflexive transitive closure of $\vdash$. Observe that a final configuration has no successor and that the initial configuration has no predecessor.

A *computation path* of length n in $M(x)$ is a string $\alpha_1, \dots, \alpha_n$ such that $\alpha_1 \vdash_{M(x)} \dots \vdash_{M(x)} \alpha_n$. We say that $M(x)$ has an *infinite computation path* if for every $n \in \mathbb{N}$ there exists a computation path of length n . A computation path $\alpha_1, \dots, \alpha_n$ is *cyclic* if $\alpha_1 = \alpha_n$.

We should be more rigorous and define the ATM to be a 6-tuple with k —the number of tapes— to be a parameter. But thanks to some classical theorems, we know that the number of tapes does not really matter, even if we have complexity concern: an Alternating Turing Machine with k tapes can reasonably be simulated with an Alternating Turing Machine with 3 tapes, one for the entry, one working tape, and one tape to write down the output.

In fact, we won't even bother taking into account the output tape in most of the case, for our Alternating Turing Machines will only *decide*, i.e. accept or reject an input, as the following definition explains.

Definition 1.1.3 (Accepted word). We define a labelling l from the configurations of $M(x)$ to $\{0, 1\}$: a configuration α of $M(x)$ is labelled with 1 iff

- $g(\alpha) = 1$, or
- $g(\alpha) = \forall$ and for every configuration β such that $\alpha \vdash \beta$, $l(\beta) = 1$, or
- $g(\alpha) = \exists$ and there exists a configuration β such that $\alpha \vdash \beta$ and $l(\beta) = 1$.

We can recursively label every configuration of $M(x)$ with l . As soon as the initial configuration of M has been labelled, we can stop the procedure and leave some configuration unlabelled.

A word $x \in \Sigma^*$ is said to be *accepted by M* if the initial configuration of $M(x)$ is labelled with 1. We set $\mathcal{L}(M) = \{x \in \Sigma^* \mid M \text{ accepts } x\}$.

Once the computation is done,⁶ that is to say when the initial configuration may be labelled, we can know what are the resources consumed by the processing of the input. We define three costs models on our ATMs.

⁶This is precisely what differentiate implicit and explicit computational complexity. Here, we have to perform the computation to know the resources needed to solve a problem.

Definition 1.1.4 (Bounds on ATM). Given M an ATM and s , t and a three functions, we say that M is space- (resp. time-, alternation-) bounded by s (resp. t , a) if for every input x of size $|x|$,

at every moment, at most $\mathcal{O}(s(|x|))$ cells of the working tapes are left non-blank,⁷

every computation path is of length at most $\mathcal{O}(t(|x|))$,

in every computation path, the number of times an existential configuration has for successor an universal configuration plus the number of times an universal configuration has for successor an existential configuration is at most $\mathcal{O}(a(|x|))$.

Note that if M is space-bounded by s then given an input x the number of configurations of M is at most $2^{\mathcal{O}(s(|x|))}$: $M(x)$ can have $\text{Card}(Q)$ different states, the working tapes can contains $k \times \text{Card}(\Sigma)^{s(|x|)}$ different words of length at most $s(|x|)$, the head on the input tape can be in $|x|$ different places and the k heads can be in $k \times s(|x|)$ different positions. To sum up, $M(x)$ can only have $\text{Card}(Q) \times (k \times \text{Card}(\Sigma)^{s(|x|)}) \times (|x|) \times (k \times s(|x|))$ different configurations. As all those value except $|x|$ are given with M , so we know there exists a d such that $M(x)$ has less than $2^{d \times s(|x|)}$ different configurations.

Moreover, if M is bounded in time by t , then M is bounded in alternation by t , for M cannot “alternate more than it computes”.

If all the configurations are existential (resp. universal), we say that M is a *non-deterministic* (resp. *universally non-deterministic*) Turing Machine. If the transition relation δ is functional, M is said to be a *deterministic* Turing Machine.

The rest of this subsection will define some other important notions like problems, reductions or complexity classes. We should dwell however on the links between Alternating Turing Machines, graphs and parallel complexity. It will be proven very useful in the rest of this thesis.

Definition 1.1.5 (Computation graph, inspired by VOLLMER [135, p. 54]). Given an ATM M and an input $x \equiv x_1 \dots x_n$, a *computation graph* $G_{M(x)}$ is a finite acyclic directed graph whose nodes are configurations of M , built as follow:

- Its roots is the initial configuration of M ,
- if $\alpha \in G_{M(x)}$ and $\alpha \vdash \beta$, then $\beta \in G_{M(x)}$ and there is an edge from α to β ,

A node in $G_{M(x)}$ is a *sink* if the associated configuration is final.

The computation graph associated to M on input x , $G_{M(x)}$, is of depth⁸ $\mathcal{O}(t(|x|))$ if M is bounded in time by t . If M is bounded in space by s , the number of configurations of M is bounded by $2^{\mathcal{O}(s(|x|))}$ and so is the number of nodes of $G_{M(x)}$. Note also that if M is deterministic, $G_{M(x)}$ is a linear graph.

To get a better insight on this construction, one may refer to Figure 1.1 where the example of a computation graph is sketched.

To determine from a computation graph if the course of computation described leads to acceptance or rejection, we simply have to label every node following the procedure described in Definition 1.1.3.

⁷The entry tape does not count in the measure of the space needed to perform a computation. We dwell on that point for it would make no sense to speak of a computation using a logarithmic amount of space in the size of the input elsewhere: one would not even have the space to write down the input! This is one of the “log-space computation oddities”, more are to come.

⁸The depth of $G_{M(x)}$ is the maximal number of edges between the node associated to the initial configuration and any other node.

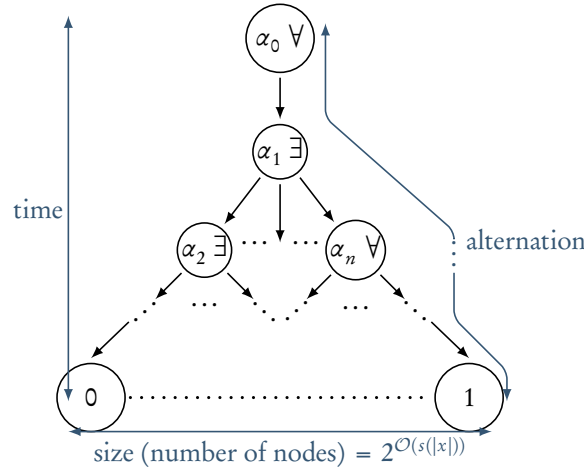


Figure 1.1: The computation graph of an ATM

In [Section 2.4](#), we will define ATMs in input normal forms ([Definition 2.4.1](#)) and use this crucial notion to define family of computation graphs ([Definition 2.4.2](#)).

This definition as the others use some classical theorems without stating them. For instance, the fact that a computational graph is *finite* and *acyclic* comes from two reasonable modifications that one can perform on any ATM. We can always add a clock to our ATMs, and forgive any transition such that $\alpha \vdash \beta \vdash^* \alpha$, i.e. to prevent any loop in our program. Notice also that we did not bother to define negative states for our ATMs, for it was proven in the very first paper on this topic [[24](#), Theorem 2.5, p. 120] that those states did not bring any computational advantage.

We will also use the following classical modifications: we can always take the alphabet Σ to be $\{0, 1\}$ and the number of working tapes k to be 1. This has some minor consequences on the complexity of our object, but stay within the realm of “reasonable” modifications, i.e. does not modify the complexity classes characterized.

Problems and complexity classes

Definition 1.1.6 ((Decision) Problem). A problem P is a subset of $\mathbb{N}$. If there exists M an ATM, such that for all $x \in \mathbb{N}$, $M(x)$ accepts iff $x \in P$, we say that M *decides* P . We will sometimes say in this case that P is the *language* of M , i.e. we define $\mathcal{L}(M) = \{n \in \mathbb{N} \mid M(x) \text{ accepts}\}$.

Definition 1.1.7 ($\text{STA}(*, *, *)$ ⁹). Given s , t and a three functions, we define the *complexity class* $\text{STA}(s, t, a)$ to be the set of problems P such that for every of them there exists an ATM bounded in space by s , in time by t and in alternation by a that decides it.

If one of the function is not specified, there is no bound on that parameter and we write $*$, and if—in the case of alternation—we allow only $\exists$ (resp. $\forall$), we write $\exists$ (resp. $\forall$) in the corresponding field. If M is deterministic, we write **d** in the alternation field.

Decisions problems are not sufficient to define everything we need, for instance because the simulation of the [Invariance Thesis](#) needs to process functions. Roughly, we say that a function f is

⁹This acronym is of course to be read as *Space, Time, Alternation*.

computable in $\mathbf{C}$ if a machine with corresponding bounds can for all $x \in \mathbb{N}$ write on a dedicated tape (“the output tape”) $f(x)$. Note that the space needed to write down this output may not be taken into account when measuring space bound, for we can consider it is a write-only tape. This is one of the reasons composition is so tricky in the log-space bounded word: if we compose two log-space bounded machines M and M' , the output-tape of M becomes one of the working tape of $M \circ M'$, and so $M \circ M'$ does not respect any more the log-space bound.

This also force us to trick the classical definition of “reductions”, that are the equivalent of simulation between machine models when applied to problems:

Definition 1.1.8 (Reduction). Given two problems P and P' , we say that P reduces to P' with a $\mathbf{C}$ -reduction and we write $P \leq_{\mathbf{C}} P'$ iff there exists a function f computable in $\mathbf{C}$ such that for all x , $x \in P$ iff $f(x) \in P'$.

There is a whole galaxy of papers studying what a reduction exactly is, for one should be really careful about the resources needed to process it. We refer to any classical textbook for an introduction to this notion, but we define anyhow log-space reduction for it will be useful.

Definition 1.1.9 (log-space reduction). A function $f : \{0, 1\}^* \rightarrow \{0, 1\}^*$ is¹⁰ said to be *implicitly log-space computable* if there exists a deterministic Turing Machine with a log-space bound that accepts $\langle x, i \rangle$ iff the i th bit of $f(x)$ is 1, and that rejects elsewhere.

A problem P is *log-space reducible* to a problem P' if there exists a function f that is implicitly log-space computable and $x \in P$ iff $f(x) \in P'$.

Reduction should always be defined accordingly to the resources consumed, but in the following we will always suppose that the “handy” reduction is taken, i.e. the one that uses the less (or supposedly less) computational power. An other crucial notion for complexity theory is the notion of *complete problem of a complexity class*: this notion somehow captures a “representative problem”, it helps to get an intuitive understanding of what a complexity class is capable of.

Definition 1.1.10 ($\mathbf{C}$ -complete problem). A problem P is *$\mathbf{C}$ -complete (under a $\mathbf{C}'$ -reduction)* if

- $P \in \mathbf{C}$
- For all $P' \in \mathbf{C}$, $P' \leq_{\mathbf{C}'} P$ with $\mathbf{C}' \subseteq \mathbf{C}$.

If P entails only the second condition, P is said to be *$\mathbf{C}$ -hard*.

We can now define some classical complexity classes, first the *sequential* complexity classes, the one that can be defined without really taking profit of alternation:

Definition 1.1.11 (Complexity classes).

$$\mathbf{L} = \mathbf{STA}(\log, *, \mathbf{d}) \tag{1.1}$$

$$\mathbf{NL} = \mathbf{STA}(\log, *, \exists) \tag{1.2}$$

$$\mathbf{co-NL} = \mathbf{STA}(\log, *, \forall) \tag{1.3}$$

$$\mathbf{P} = \mathbf{STA}(\log, *, *) = \mathbf{STA}(*, \text{poly}, \mathbf{d}) \tag{1.4}$$

¹⁰Some hypothesis are required, we refer for instance to [4, Chapter 4: Space complexity] for a precise statement.

We introduced here the *complementary* of a complexity class, namely **co-NL**. Let P be a problem, $PComp$ is $\{n \in \mathbb{N} \mid n \notin P\}$. For any complexity class C , **co-C** denotes the class $\{PComp \mid P \in C\}$. If a problem $P \in C$ can be decided with a deterministic machine, it is trivial that $PComp$ can be decided with the same resources, it is sufficient to inverse acceptance and rejection, and so $C = \mathbf{co-C}$.

We chose here to *define* the classical parallel classes of complexity **AC** and **NC** in terms of Alternating Turing Machines, the other definition in terms of *Boolean circuits* being postponed to Definition 2.1.5. The correspondence between those two definitions is a mix of classical results [15, p. 14] (eq. (1.6)), [25], [123], [115] (eq. (1.8)).

Definition 1.1.12 (Alternation and parallel complexity classes). For all $i \geq 1$

$$AC^0 = STA(*, \log, 1) \quad (1.5)$$

$$NC^1 = STA(*, \log, *) \quad (1.6)$$

$$AC^i = STA(\log, *, \log^i) \quad (1.7)$$

$$NC^i = STA(\log, \log^i, *) \quad (1.8)$$

We define $\cup_i AC^i = AC$ and $\cup_i NC^i = NC$.

Theorem 1.1.1 (Hierarchy).

$$AC^0 \subsetneq NC^1 \subseteq L \subseteq NL = \mathbf{co-NL} \subseteq AC^1 \subseteq NC^2 \subseteq \dots \subseteq AC = NC \subseteq P$$

For all $i \in \mathbb{N}$, $AC^i \subseteq NC^{i+1}$ is quite evident, for a bound of $\log^{i+1}$ on time is sufficient to perform at least $\log^i$ alternations. That $AC^0 \subsetneq NC^1$ is a major result of FURST, SAXE, and SIPSER [44] and made the Boolean circuits gain a lot of attractiveness, for separation results are quite rare!

The proof of $NL \subseteq NC^2$ is rather nice, and one may easily find a clear explanation of it [135, p. 41, Theorem 2.9]: a matrix whose dimension is the number of configurations of a log-space Turing Machine is built in constant depth,¹¹ and a circuit of depth $\log^2$ computes its transitive closure in parallel.

At last, $\mathbf{co-NL} = NL$ is a major result of IMMERMANN [78]. It was proven by showing how a log-space bounded Turing Machine could build the computation graph of another log-space bounded Turing Machine and parse it to determine if there is a path from the initial configuration to an accepting configuration. That $\mathbf{co-L} = L$ is trivial, since every deterministic complexity class is closed by complementation.

Turing Machines indeed share a lot with graphs, especially when they are log-space bounded. We illustrate that point with two (decision) problems that we will use later on. They both are linked to connectivity in a graph, and can be also refereed to as “reachability” or “maze-problem”. Numerous papers are devoted to list the problems complete for a class, we can cite among them JENNER, LANGE, and MCKENZIE [81] and JONES, LIEN, and LAASER [85].

Problem A: *Undirected Source-Target Connectivity (USTConn)*

Input: An undirected graph $G = (V, E)$ and two vertices $s, t \in E$

Output: Yes if s and t in the same connected component in G , i.e. if there is a path from s to t .

¹¹That is to say with a family of constant-depth unbounded Boolean circuit. It might just be understand for the time being as “with AC^0 resources”, a more precise definition will be given in Chapter 2.

REINGOLD [114] proved that *USTConn* is L-complete (hardness was proved by COOK and MCKENZIE [28]). We will use in Chapter 2 its bounded variant, where G is of degree at most n , *USTConn*₂. We will also use in Appendix B a small variant of its complementary, where s and t are not parameters, but always the first and the last nodes of V .

This problem also exists in a directed fashion:

Problem B: (*Directed*) *Source-Target Connectivity* (*STConn*)

Input: A directed graph $G = (V, A)$ and two arcs $s, t \in A$

Output: Yes if there a path from s to t in G .

STConn is NL-complete.

The proof that every NL-problem can be reduced to *STConn* is really interesting, for it links closely log-space computation and graphs.¹² Let $P \in \text{NL}$, there exists non-deterministic log-space bounded Turing Machine M that decides it (otherwise there would be no evidence that $P \in \text{NL}$). There exists a constant-depth boolean circuit that given $n \in \mathbb{N}$ and the description of M outputs $G_{M(n)}$ the configuration graph of M . Then $n \in P$ iff there is a path between the initial configuration and a configuration accepting in $G_{M(n)}$, a problem that is *STConn*.

1.2 Linear Logic

Linear Logic (LL) was invented by GIRARD[47] as a refinement of classical logic. Indeed, (classical) logic is a pertinent tool to formalise “the science of mathematical reasoning”, but when we affirm that a property A entails a property B ($A \Rightarrow B$), we do not worry about *the number of times A is needed to prove B* . In other words, there is no need to take care of how resources are handled, since we are working with ideas that may be reproduced for free.

We just saw that if we have complexity concerns, one should pay attention to duplication and erasure, for it has a cost. The implication $\Rightarrow$ superimposes in fact two operations that Linear Logic distinguishes: first, there is the *linear implication* $A \multimap B$ that has the meaning of “ A has to be consumed once to produce B ”, and the *of course modality* $!A$ that has the meaning of “I can have as many copies (including 0) of A as I want”. Thus, the classical implication $A \Rightarrow B$ is encoded as $(!A) \multimap B$. From this first remark, numerous tracks were followed: we will begin by recalling the main ideas of Linear Logic, and then focus on proof nets before presenting the Geometry of Interaction (GoI) program.

The reference to study Linear Logic remains *Proofs and Types*[62], and the reader not used to this subject could read with benefits the 10 pages of Appendix B, “What is Linear Logic?”, written by LAFONT. For anyone who reads french, JOINET’s *Habilitation à diriger les recherches*[83] is an excellent starting point to put Linear Logic into philosophical and historical perspective. The two volumes of GIRARD’s textbook [55, 56] were translated in English[58] and are the most complete technical introduction to the subject.

Main concepts and rules

We will tend during the course of this work to dwell on the syntactical properties of Linear Logic, but we should however remark that some of its elements appeared first in LAMBEK [94]’s linguistic considerations. Linear Logic is a world where we cannot duplicate, contract and erase occurrences

¹²For a more advanced study of those links, one could have a look at the work led by WIGDERSON [136] and to its important bibliography.

$$\begin{array}{c}
\frac{}{\vdash A, A^\perp} ax. \quad \frac{\vdash \Gamma, A \quad \vdash \Delta, B}{\vdash \Gamma, \Delta, A \otimes B} \otimes \quad \frac{\vdash \Gamma, A}{\vdash \Gamma, ?A} der. \quad \frac{\vdash \Gamma, A}{\vdash \Gamma, !A} ! \\
\frac{\vdash \Gamma, A \quad \vdash \Delta, A^\perp}{\vdash \Gamma, \Delta} cut \quad \frac{\vdash \Gamma, A, B}{\vdash \Gamma, A \wp B} \wp \quad \frac{\vdash \Gamma}{\vdash \Gamma, ?A} ? \quad \frac{\vdash \Gamma, ?A, ?A}{\vdash \Gamma, ?A} ctr.
\end{array}$$

From left to right, the names and the group of the rules are

Group	Structural	Multiplicative	Exponential
1st line	Axiom	Tensor introduction	Dereliction Promotion
2nd line	Cut	Par introduction	Weakening Contraction

Figure 1.2: The rules of MELL

of formula for free. This led to the division of classical logic¹³ into three *fragments*: additive, whose rules are context-sensitive, multiplicative, whose rules are reversible, and exponential, which handles the modalities. The negation of a formula is encoded thanks to the *linear negation* $^\perp$, which is an involution.

We will only use the MELL fragment (without units),¹⁴ i.e. the multiplicative and exponential connectives, whose formula are, for α a *literal* (or *atom*, a variable taken for a countable set), using the inductive *Backus Normal Form* notation:

$$A := \alpha \mid \alpha^\perp \mid A \otimes A \mid A \wp A \mid !A \mid ?A.$$

The $\otimes$ connective is called *tensor*, the $\wp$ connective is called *par*, $!$ and $?$ are respectively named of *course* (or *bang*) and *why not*. Dualities are defined with respect to De Morgan's laws, that allows us to let the negation flows down to the literals:

$$(A \otimes B)^\perp := A^\perp \wp B^\perp \quad (A \wp B)^\perp := A^\perp \otimes B^\perp \quad (!A)^\perp := ?(A^\perp) \quad (\alpha^\perp)^\perp := \alpha$$

The connective $\multimap$ is defined as $A \multimap B \equiv (A^\perp) \wp B$. We let Γ and Δ be two multisets of formulae, and for $\diamond \in \{!, ?\}$, if $\Delta = A_1, \dots, A_n$, we let $\diamond \Delta$ be $\diamond A_1, \dots, \diamond A_n$.

We borrow from MOGBIL [106, p. 8] the following presentation of rules and proofs, which is especially handy. A *sequent* $\vdash \Gamma$ is a multiset of formulae of MELL, the order of the occurrences of formula plays no role for the time being. We let S be the set of sequents of MELL. For $n \geq 0$, a n -ary *rule* R_n with n hypothesis and one conclusion is a $n + 1$ relation on S , i.e. $R \subseteq S^n \times S$. The set of rules of MELL is given in Figure 1.2, where the *principal formula* is distinguished in the conclusion sequent.

A *proof of height 0* in MELL is the unary rule *ax.*. For $n \neq 0$, a *proof of height $h \neq 0$* D_h is a subset of $(D_{<h})^n \times R_n$ such that

- the set $D_{<h}$ is the set of proofs of height strictly inferior to h ,
- the n conclusions of the n proofs of $D_{<h}$ are the n premises of R_n ,
- the conclusion of D_h is the conclusion of R_n .

¹³For a complete presentation of the variations of Linear, Classical and Intuitionistic Logic, we refer for instance to LAURENT [95].

¹⁴In a nutshell, it is natural to associate to the multiplicative fragment a parallel handling of the resources, whereas the additive fragment is rather to be conceived as highly non-deterministic.

A sequent $\vdash \Gamma$ is derivable in MELL if it admits a proof, i.e. if there exists a proof π whose conclusion is $\vdash \Gamma$. The η -expansion of a proof π is a proof π' of same conclusion where the axioms introduce only literals. Such a proof π' always exists, and it is unique. A proof π is α -equivalent¹⁵ to a proof π' if π' can be obtained from π by replacing the occurrences of the literals in π .

One of the fundamental theorem of logic is GENTZEN's *Hauptsatz*, or cut-elimination theorem, proved first by GENTZEN [45]¹⁶ for classical and intuitionistic logic. This theorem applies as well to MELL, and we sometimes say that a system *eliminates cuts* if it entails this theorem:

Cut-elimination Theorem. *For every proof π of $\vdash \Gamma$ in MELL, there exists a proof π' of same conclusion that does not uses the rule cut.*

The Curry-Howard-correspondence (*a.k.a.* the *proof-as-program* correspondence) tells us that the operation of rewriting a proof with cuts into a proof without cuts corresponds to the execution of a program written in λ -calculus. This link between computer programs and mathematical proofs justifies the use of Linear Logic as a typing system on programs, and by extension, its use as a tool to modelise complexity. Before studying in depth how this gateway was used, we still need to introduce two other key-concepts of Linear Logic: proof nets and Geometry of Interaction (GoI). They correspond respectively to the model we develop in Chapter 2, and to the model we develop in Chapter 3.

Proof Nets

Proof Nets are a parallel syntax to express proofs that get rid of some lumbering order in the application of rules. For instance, one should not feel the urge to distinguish between the following two proofs obtained from the same proof π :

$$\begin{array}{c} \vdots \pi \\ \hline \frac{\vdash A, B, C, D}{\vdash A \wp B, C, D} \wp \\ \hline \vdash A \wp B, C \wp D \end{array} \quad \begin{array}{c} \vdots \pi \\ \hline \frac{\vdash A, B, C, D}{\vdash A, B, C \wp D} \wp \\ \hline \vdash A \wp B, C \wp D \end{array}$$

It was also introduced by GIRARD[51] and provoked numerous studies, and yet there is to our knowledge no “classical definition”. The most accurate one can be found in MOGBIL's *Habilitation à diriger les recherches* [106, Definition 2.11, p. 11], it defines proof nets as hyper-multigraphs.

We will present in details and massively use proof nets for an unbounded version of Multiplicative Linear Logic in Chapter 2. So despite the tradition to present some drawing to introduce them, we will just express some global ideas and refer the reader to Section 2.1 for a complete presentation. We will use them only as examples and hints toward a sharper understanding in Chapter 3, and they do not appear in Chapter 4.

A proof net is built inductively by associating to every rule of Figure 1.2 a rule on drawings. We obtain by doing so a labeled graph, where the nodes are endowed with a connective of MELL, and the edges represents sonships between occurrences of formulae. This parallel writing of the proofs gives a quotient on proofs that “reduces the syntactical noise” and by doing so let the proofs be more readable. Every derivation infers a proof net, and by applying some quotient on the different representations of the same proof net, a proof net may be inferred by several proofs.

¹⁵We won't define this notion precisely here, for it is very tedious to do so, as attested by the 1st chapter of KRIVINE [92].

¹⁶The interested reader may find a translation [46].

One should remark that none of those rules is local in this setting, for there is always (except for the axiom rule) a constraint on the membership of the principal occurrences to the same context or to two different contexts, not to state the constraint on the presence of a modality. This is the drawback of this representation: the correctness of a proof is not tested locally any more, as opposed to any other syntactical system that existed in the history. We call *pseudo nets* the structures freely-generated, that is, such that we don't know if there exists a proof that infers it. To determine whether a pseudo net is a proof net, we need a global criteria to test its "legacy". Those criterion are called criterion of correctness and we refer once again to MOGBIL's work [106, p. 16] for a complete list. We will use one of them in [Theorem 2.2.1](#).

We left aside the presentation of those (yet very important) tools to focus on a crucial point for the well understanding of this thesis. As we presented it, the sequents of MLL are not ordered. An order could be introduced thanks to a bit tedious notion of *exchanges* of occurrences of formulae, allowing us to keep track of their positions. It is handy not to take care of this aspect, but we are forced to re-introduce it when dealing with proof nets: suppose given a proof π of conclusion $\vdash ?A, ?A, ?A, B, B$. The proof π can be completed as follows:

$$\frac{\vdash \pi \quad \frac{\vdash ?A, ?A, ?A, B, B}{\vdash ?A, ?A, B, B} \text{ctr.}}{\vdash ?A, ?A, B \wp B} \wp$$

The question is: if we label the conclusion of π with $\vdash ?A_1, ?A_2, ?A_3, B_1, B_2$, what is the conclusion of the previous proof? It could be any sequent among

$$\begin{array}{lll} \vdash ?A_1, ?A_2, B_1 \wp B_2 & \vdash ?A_1, ?A_3, B_1 \wp B_2 & \vdash ?A_2, ?A_3, B_1 \wp B_2 \\ \vdash ?A_1, ?A_2, B_2 \wp B_1 & \vdash ?A_1, ?A_3, B_2 \wp B_1 & \vdash ?A_2, ?A_3, B_2 \wp B_1 \end{array}$$

In fact, by dealing with non-ordered sequents, we already introduced a first quotient on our proofs, which is not compatible with proof nets. Remark that it does not contradict the fact that proof nets quotient proofs, for we may for this example performs the contractions and the par-introduction in parallel.

But in the following, we will encode the truth-value of a boolean type in the order of the the premise of a par-rule ([Definition 2.2.1](#)). The principal formulae of the contractions will later on encode the value of a bit in a string ([Section 3.3](#)). This distinction, which by the way appears in Geometry of Interaction, is crucial since it may be the only difference between the proofs corresponding to two different binary lists.

Getting closer to a Geometry of Interaction

Geometry of Interaction (GoI) tended to be used as a motto to qualify several approaches to reconstruct logic. As JOINET [83, p. 43] points it, the heading "Geometry of Interaction" is used at the same time to describe

- a way to rebuild logic from computational interaction,
- a formalization of the preceding idea in a functional analysis framework.

Their common point is a strong will to “rebuild logic from the inside”, by letting the cut-elimination—and by extension the dynamic—be the primary notion from which logic is built, and hence to obtain immediately computational properties. This subsection does not recollect all the previous attempts to work in the GoI setting, but gives some hints of its goal, its links with complexity being postponed to the next section and to [Section 3.4](#).

The name GoI was first introduced by GIRARD [49] and it was accompanied with a clear justification that we cite entirely:

“The only extant semantics for computation are denotational, i.e. static. This is the case for the original semantics of Scott [118], which dates back to 1969, and this remains true for the more recent coherent semantics of the author [47]. These semantics interpret proofs as functions, instead of actions. But computation is a dynamic process, analogous to —say— mechanics. The denotational approach to computation is to computer science what statics is to mechanics: a small part of the subject, but a relevant one. The fact that denotational semantics is kept constant during a computational process should be compared to the existence of static invariants like mass in classical mechanics. But the core of mechanics is dynamics, where other invariants of a dynamical nature, like energy, impulsion etc. play a prominent role. Trying to modelize programs as actions is therefore trying to fill the most obvious gap in the theory. There is no appropriate extant name for what we are aiming at: the name ‘operational semantics’ has been already widely used to speak of step-by step paraphrases of computational processes, while we are clearly aiming at a less ad hoc description. This is why we propose the name

geometry of interaction

for such a thing.”

(GIRARD [49, p. 26])

Starting from the consideration that “linear negation behaves like transposition in Linear Algebra” [49, p. 10], a community started to work on the building of the Logic in a Linear Algebra framework. One leading idea was to define *ex nihilo* the concept of type, and to replace cut-elimination by the action of some operators on the Hilbert space. Proofs were interpreted as symmetrical matrices, i.e. equal to their own transpositions, and they were, so to say, “the written trace of underlying geometrical structures” [49, p. 1]. What was really interesting was the ability to express strong normalisation as nilpotency thanks to the “Execution formula”. GoI was naturally conceived as a “new form of semantics” and expressed within a C^* -algebra.

Yet, coherent and Banach spaces had their limitations:

“quantum coherent spaces [...] have [...] two major drawbacks: that of a categorical interpretation, whence unable to explain dynamics, thus complexity; and, at a deeper level, their incompatibility with infinite dimension, the latter drawback being related to the former.”

(GIRARD [57, p. 26])

So the von Neumann algebra framework was developed, but we save this part for [Chapter 3](#). One could however mention that this approach is a more general theory, for the subtle notions of finiteness helps to understand infinity in logic and in complexity.

A concrete version of the interpretation of cut-elimination by actions was provided by the context semantics of proof nets. This could be interpreted as the wandering of token, starting from the conclusions of the proof net and eventually exiting by some conclusion of the graph. Those pebbles moving in the structure without modifying it led to the study of *paths* in proof net [5], the normal form of the proof net being given by *persistent paths*. This led to the study of several kinds of *token machine* [96] and to the development of the study of algebra weight.

1.3 Previous Gateways between Linear Logic and Complexity

We will rapidly scan the wide spectrum of links between Linear Logic and complexity. This small section has two purposes: to present the numerous previous approaches, and that should ease the understanding of our contribution. We won't precise the settings that were developed, but gives several key-references for the reader who would like to push this study further.

First of all, Linear Logic and complexity are linked because some problems regarding Linear Logic arise naturally, and we can try to establish the complexity of those problems. This approach led to numerous interesting problems, and allowed a better understanding of the nature of Linear Logic.

Cut-elimination Given two proofs and a confluent cut-elimination procedure, do they have the same cut-free form? MAIRSON and TERUI [100] established a hierarchy, proving that this problem is **P**-complete for MLL and **co-NP**-complete for MALL.

Correctness Given a proof-structure, is it a proof net? This problem was proven to be **NL**-complete no matter what the fragment considered is by JACOBÉ DE NAUROIS and MOGBIL [80].

Provability Given a formula, is it provable? This problem was proved to be **NP**-complete for MLL by KANOVICH [87]. The reader may also have a look at LINCOLN [98] and its consequent bibliography to learn more about the complexity of theorem proving.

Those problems give insights on the complexity of Linear Logic, but other approaches have been developed to address this question. One can use Linear Logic to deliver certification to programs, attesting that they will “behave well” no matter the context. This needs to define subsystems of LL and to establish that any program typable with this subsystem will be of such or such complexity.

To get rid of a *machine-based* characterisation of complexity classes, one could also perform the opposite operation: to design a “logic language of programming”. We get sure that any program written in this syntax will be “pre-constrained”

The constraint on Linear Logic one could develop to moderate its expressivity was naturally the modalities. One of the oldest—and yet, very accurate—formulation of this idea is the following:

“[T]he rule for ‘!’ indicates unlimited possibilities of duplication, but not a concrete one [...]. The clarification of this point could be of great importance: consider for instance bounded exponentials $! \alpha A$, $? \alpha A$, that could be added to linear logic with the intuitive meaning of ‘iterate α times’. [T]here is some underlying polynomial structure in the exponentials. Now, it is not always the case that we can associate polynomials to all exponentials occurring in a proof of standard linear logic, especially when we have to deal with cut ; hence the proofs admitting such polynomial indexings are very peculiar.”

(GIRARD [49, p. 93])

This idea was later on developed by GIRARD, SCEDROV, and SCOTT [63] under the name of “Bounded Linear Logic”. It was later on refined [52], made more precise by DANOS and JOINET [37] with the introduction of *Elementary Linear Logic* (ELL) and generalized by the introduction of *Linear Logic by level* by BAILLOT and MAZZA [13]. Those works managed to characterise elementary and deterministic poly-time computation in Linear Logic, using sequents and proof nets. SCHÖPP [117] adapted this idea to second-order quantifier with benefits, for it gave a characterization of **L**.

Numerous studies were also developed with a purely semantical approach, but there is a conflict between their will to “characterise all the programs that compute the same function” and our objective to “characterise *the* good program that computes a function”. The semantics is somehow “talkative”, whereas the syntax focuses on the good programs, the more significant. A quote—with our emphasis—may clarify this point:

“A different approach to the semantics of bounded time complexity is possible: the basic idea is to measure by semantic means the execution of any program, regardless to its computational complexity. The aim is to compare different computational behaviours and to learn afterwards something on the very nature of bounded time complexity.”

(DE CARVALHO, PAGANI, and TORTORA DE FALCO [39])

One should remark moreover that this setting is not an Implicit Complexity approach at all.

Few works has been led in the GoI setting, which should recollect the advantages of the syntax and the advantages of semantics. Among them, it is worth referring to BAILLOT and PEDICINI [14], which proposed an interpretation of Elementary Linear Logic by an algebra of clause equipped with resolution. We can also, to conclude this small scanning of the previous approaches, mention the work led by DAL LAGO and SCHÖPP [35]. It is relevant for it introduces the functional programming language INTML that rely on Linear Logic. What is interesting is that it led to an actual implementation, whose source may be found at <https://github.com/uelis/IntML>.

In the Following...

The “level-by-level” approach has now reached a mature state. There is of course some enhancement that could be developed, but this structure begins to be quite well understood, capturing more and more programs, being less constraining. The work led by SCHÖPP [117] and INTML [35] are to our knowledge the only attempts to capture classes below **P**.

The two approach we are going to analyse do not rely heavily on the modalities to constrain the power of Linear Logic. In fact, they do not even appears in Boolean proof nets, which handle the computation by a kind of “hard-wire” setting: everything will be handled by the crossing of edges, or... wires. This will take benefit from the parallel writing of proofs as proof nets, but we still have to define a parallel execution for them. The next chapter will expose this setting and link it both to Boolean circuits and Alternating Turing Machines.

The other approach will rely on the GoI program, for we will encode proofs of integers in a von Neumann algebra and interpret nilpotency as normalisation, or in this case, acceptance. It will be wildly non-deterministic, for the computation will take place in product of matrices computing in parallel. Surprisingly, this framework can be “decided” with log-space resources, and that will conclude our Chapter 3.

To prove that we capture log-space computation, we will use in [Chapter 4](#) abstract machines that manipulates pointers. This is a long-standing goal, as attested by a quotation that will conclude this introductory chapter:

“[G]etting rid of syntax [...] is the ultimate aim, which has been achieved only for multiplicatives [...] by Vincent DANOS and Laurent REGNIER [38]. In this case, the basic underlying structure turns out to be permutation of pointers, and the notion of a cyclic permutation plays a prominent role (DANOS and REGNIER replaced cyclicity by conditions involving connected acyclic graphs, i.e. trees). It seems that the general solution (not yet known, even for quantifiers) could be something like permutations of pointers, with variable addresses, so that some (simple form of) unification could be needed for the composition of moves : yet another connection with logic programming !”

(GIRARD [49, p. 101])

Proof circuits

Contents

1.1	Computational Complexity Theory	4
	Complexity and models of computation	
	Alternating Turing Machine	
	Problems and complexity classes	
1.2	Linear Logic	12
	Main concepts and rules	
	Proof Nets	
	Getting closer to a Geometry of Interaction	
1.3	Previous Gateways between Linear Logic and Complexity	17

PROOF nets are the traditional way of representing proofs of Linear Logic in parallel. And yet, they were not until the work by TERUI[130] really used to model parallel computation. However, it is rather easy to define a parallel cut-elimination procedure and to extend the Proof-as-program correspondence to parallel evaluation in this setting. The ideal target to this extension are Boolean circuits, which are the classical parallel model of computation: they can be seen as electronic or digital circuits made of logic gates computing simple functions.

One of the computational property of Boolean circuits is that they are able to duplicate temporary values previously computed and to distribute them to several logic gates. We saw that duplication could be handled in an elegant way in Linear Logic by the ! and ? modalities, and yet we won't use this tool in the setting we propose here. Instead, we will build proof nets as hardware devices, where everything is explicit and where the duplication is "pre-constrained".

The model of computation we are going to define is close to graph rewriting, where the proof nets are made of only the multiplicative fragment of Linear Logic, and so every operation will be local. Because of the linearity of this logic, we are forced to keep track of the partial results generated by the evaluation that are unused in the result. It is a low-level of computation and every algorithm that can be written on Boolean circuits will be "in hard" translated into Boolean proof nets.

Regarding this first link, we will use the approach developed by MOGBIL and RAHLI [107, 105], who focused on uniform classes of complexity. We will present a novel translation in AC^0 from Boolean proof nets to Boolean circuits by focusing on a simpler restricted notion of uniform Boolean proof nets—proof circuits. That will allow us to compare complexity classes below log-s pace, which were out of reach with the previous translations.

The second bridge we are going to build is between proof circuits and Alternating Turing Machines (ATMs). It will help us to present ATMs as a convenient model of parallel computation, and to prove that in some cases, logarithmic space is enough to establish the normal form of a proof net (see Section 2.4). Last but not least, this results exhibits a difference between our unbounded proof circuits and the bounded ones, something that could not be seen with the mere Boolean proof nets.

Regarding parallel complexity, a gentle introduction (in french) to this subject may be found in *Les petits cailloux* [112], but the most complete handbooks remains *Introduction to Circuit Complexity: A Uniform Approach* [135] and *Limits to Parallel Computation : P-Completeness theory* [65].

2.1 Basic Definitions: Boolean Circuits and MLL_u

We begin by introducing Boolean circuits: knowing them is not mandatory to understand how MLL_u proof nets compute, but it is useful to bear in mind notions like *uniformity*, *families of Boolean circuits* or *depth* of a Boolean circuit. Proof nets for MLL_u are an efficient way to model parallel computation with Linear Logic.

Boolean circuits

Boolean circuits (Definition 2.1.2) were first defined and popularized as a parallel model of computation by BORODIN [19]. One of their feature is that they work only on inputs of fixed length, and that forces us to deal with *families* of Boolean circuits—and there arises the question of *uniformity* (Definition 2.1.4).

Definition 2.1.1 (Boolean function). A n -ary Boolean function f^n is a map from $\{0, 1\}^n$ to $\{0, 1\}$. A Boolean function family is a sequence $f = (f^n)_{n \in \mathbb{N}}$ and a basis is a set of Boolean functions and Boolean function families. We define the two classical basis¹:

$$\mathfrak{B}_b = \{\mathcal{N}_{eg}, \vee^2, \wedge^2\} \text{ and } \mathfrak{B}_u = \{\mathcal{N}_{eg}, (\vee^n)_{n \geq 2}, (\wedge^n)_{n \geq 2}\}$$

The Boolean function $USTConn_2$ solves the problem $USTConn_2$ (Problem A), i.e. given in input the coding of an undirected graph G of degree at most 2 and two names of gates s and t , this function outputs 1 iff there is a path between s and t in G .

The purpose of this function will be detailed later on, and illustrated in Figure 2.4.

Definition 2.1.2 (Boolean circuit). Given a basis $\mathfrak{B}$, a Boolean circuit over $\mathfrak{B}$ with n inputs C is a directed acyclic finite and labelled graph. The nodes of fan-in 0 are called *input nodes* and are labelled with $x_1, \dots, x_n, 0, 1$. Non-input nodes are called *gates* and each one of them is labelled with a Boolean function from $\mathfrak{B}$ whose arity coincides with the fan-in of the gate. There is a unique node of fan-out 0 which is the *output gate*.² We indicate with a subscript the number of inputs: a Boolean circuit C with n inputs will be named C_n .

¹Sometimes denoted respectively $\mathfrak{B}_0$ and $\mathfrak{B}_1$, but we adopt here the initials of *bounded* and *unbounded* as a reminder.

²We take into account in this definition only *decision circuits*, that is circuits that compute the characteristic function of a language. We will later on compose Boolean circuits with more than one output but took the liberty not to define them formally.

The *depth* of a Boolean circuit C_n , denoted by $d(C_n)$, is the length of the longest path between an input node and the output gate. Its *size* $|C_n|$ is its number of nodes. We will only consider Boolean circuits of size $n^{\mathcal{O}(1)}$, that is to say polynomial in the size of their input.

We say that C_n *accepts* (resp. *rejects*) a word $w \equiv w_1 \dots w_n \in \{0, 1\}^n$ iff C_n evaluates to 1 (resp. 0) when $w_1, \dots, w_n$ are respectively assigned to $x_1, \dots, x_n$. A *family of Boolean circuits* is an infinite sequence $C = (C_n)_{n \in \mathbb{N}}$ of Boolean circuits, C *accepts a language* $X \subseteq \{0, 1\}^*$ iff for all $w \in X$, $C_{|w|}$ accepts w , and for all $w' \notin X$, $C_{|w'|}$ rejects w' .

Boolean circuits are “the” classical model of parallel computing, because we consider that their evaluation is linear in their depth: every gate is an independent processor evaluating its inputs as they are “ready”. All the gates at the same depth are taken to compute at the same time, so depth is often compared to time on a sequential model of computation: the depth reflects the number of intermediate processing needed to output a value.

We now recall the definition of the *Direct Connection Language* of a family of Boolean circuits, an infinite sequence of tuples that describes it totally.

Definition 2.1.3 (Direct Connection Language RUZZO [115, p. 371]). Given $(\cdot)$ a suitable coding of integers and C a family of Boolean circuits over a basis $\mathfrak{B}$, its *Direct Connection Language* $L_{DC}(C)$ is the set of tuples $\langle 1^y, \bar{g}, \bar{p}, \bar{b} \rangle$, such that g is a gate in C_y , labelled with $b \in \mathfrak{B}$ if $p = \epsilon$, else b is its p th predecessor.

We use the unary encoding to stay within suitable bounds for the uniformity, whose definition follows.

Definition 2.1.4 (Uniformity BARRINGTON, IMMERMAN, and STRAUBING [15]). A family C is said to be *DLOGTIME*-uniform if there exists a deterministic Turing Machine with random access to the input tape³ that given $L_{DC}(C)$, 1^n and $\bar{g}$ outputs in time $\mathcal{O}(\log(|C_n|))$ any information (position, label or predecessors) about the gate g in C_n .

We can adopt the fine tuning of *DLOGTIME* many-one reductions developed by REGAN and VOLLMER [113]. This notion preserves transitivity and the membership to individual levels of hierarchy. Despite the fact that a *DLOGTIME* Turing Machine has more computational power than a constant-depth circuit, “a consensus has developed among researchers in circuit complexity that this *DLOGTIME* uniformity is the ‘right’ uniformity condition” (HESSE, ALLENDER, and BARRINGTON [70, p. 698]) for small complexity classes. Some proposal towards a circuit uniformity “sharper” than *DLOGTIME* were made by BONFANTE and MOGBIL [18], but we will stick to this well-accepted notion of uniformity and any further reference to uniformity in this chapter is to be read as *DLOGTIME* uniformity.

Uniformity allows us to be sure that there is a “sufficiently simple” way of describing the conception of any member of the family of Boolean circuits. Without it, even the problem not computable becomes computable, simply define a family of Boolean circuits that solves this problem: as you don’t have to describe the way they are built —what uniformity is—, you don’t have to bother if they may be constructed or not! Defining non-uniform classes of complexity led to some interesting questions and results, but it won’t be the topic here, for we focus on a more tractable approach of algorithms.

³Meaning that the Turing Machine, given $n \in \mathbb{N}$, can in one step read the value written in the n th cell. For some fine tuning of this notion, we refer to CAI and CHEN [23].

Definition 2.1.5 ($\mathbf{AC}^i$, $\mathbf{NC}^i$). For all $i \in \mathbb{N}$, given $\mathfrak{B}$ a basis, a language $X \subseteq \{0, 1\}^*$ belongs to the class $\mathbf{AC}^i(\mathfrak{B})$ (resp. $\mathbf{NC}^i(\mathfrak{B})$) if X is accepted by a uniform family of polynomial-size, $\log^i$ -depth Boolean circuits over $\mathfrak{B}_u \cup \mathfrak{B}$ (resp. $\mathfrak{B}_b \cup \mathfrak{B}$). We set $\mathbf{AC}^i(\emptyset) = \mathbf{AC}^i$ and $\mathbf{NC}^i(\emptyset) = \mathbf{NC}^i$.

Of course this definition of $\mathbf{AC}^i$ and $\mathbf{NC}^i$ is equivalent to the definition *via* Alternating Turing Machines we gave previously in [Definition 1.1.12](#).

Unbounded Multiplicative Linear Logic

We will here briefly expose the Logic underneath our proof circuits, a fragment of Linear Logic without modalities. We introduce those notions rapidly, for they were already developed in [Section 1.2](#) and are quite classical.

Rather than using Multiplicative Linear Logic (MLL), we work with *unbounded*⁴ *Multiplicative Linear Logic* ($\mathbf{MLL}_u$) which differs only on the arities of the connectives but better relates to circuits (see the remark about “unfolding $\mathbf{MLL}_u$ ” for a better insight). The following definitions are partially *folklore* since the introduction of “general multiplicative rule” by DANOS and REGNIER [38].

Definition 2.1.6 (Formulae of $\mathbf{MLL}_u$). Given α a literal and $n \geq 2$, formulae of $\mathbf{MLL}_u$ are:

$$A := \alpha \mid \alpha^\perp \mid \otimes^n(A_1, \dots, A_n) \mid \wp^n(A_n, \dots, A_1)$$

We write $\overrightarrow{A}$ (resp. $\overleftarrow{A}$) for an ordered sequence of formulae $A_1, \dots, A_n$, (resp. $A_n, \dots, A_1$). Duality is defined with respect to De Morgan’s law:

$$(A^\perp)^\perp := A \qquad (\otimes^n(\overrightarrow{A}))^\perp := \wp^n(\overleftarrow{A^\perp}) \qquad (\wp^n(\overleftarrow{A}))^\perp := \otimes^n(\overrightarrow{A^\perp})$$

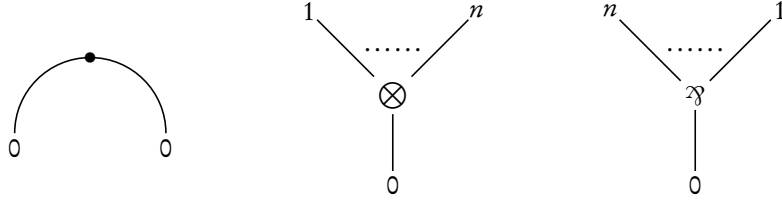
As for the rest of this chapter, consider that A , B and D will refer to $\mathbf{MLL}_u$ formulae. The result of the substitution of all the occurrences of B by an occurrence of D in A will be written $A[B/D]$, $A[D]$ if $B = \alpha$. We will sometimes write $A \wp B$ (resp. $A \otimes B$) for $\wp^2(A, B)$ (resp. $\otimes^2(A, B)$).

Definition 2.1.7 (Sequent calculus for $\mathbf{MLL}_u$). A *sequent* of $\mathbf{MLL}_u$ is of the form $\vdash \Gamma$, where Γ is a multiset of formulae. The *inference rules* of $\mathbf{MLL}_u$ are as follow:

$$\begin{array}{c} \frac{}{\vdash A, A^\perp} ax. \qquad \frac{\vdash \Gamma_1, A_1 \quad \dots \quad \vdash \Gamma_n, A_n}{\vdash \Gamma_1, \dots, \Gamma_n, \otimes^n(\overrightarrow{A})} \otimes^n \\[2ex] \frac{\vdash \Gamma, A \quad \vdash \Delta, A^\perp}{\vdash \Gamma, \Delta} cut \qquad \frac{\vdash \Gamma, \overleftarrow{A}}{\vdash \Gamma, \wp^n(\overleftarrow{A})} \wp^n \end{array}$$

Proofs of $\mathbf{MLL}_u$ are built with respect to those rules, following the construction explained page 13. $\mathbf{MLL}_u$ has neither weakening nor contraction, but admits implicit exchange and eliminates cuts. The formulae A and $A^\perp$ in the rule *cut* are called the *cut formulae*.

⁴That has nothing to do with Bounded Linear Logic: both have complexity concern, it is just that in our case we take *the connectives* to be unbounded, and not *the modalities*, that are by nature unbounded.

Figure 2.1: The components of a proof net: ax -link, $\otimes^n$ -link and $\mathfrak{M}^n$ -link

Proof Nets for MLL_u

Proof nets are a parallel syntax for MLL_u that abstracts away everything irrelevant and only keeps the structure of the proofs. We introduce measures (Definition 2.1.10) on them in order to study their structure and complexity, and a parallel elimination of their cuts (Definition 2.1.11).

Definition 2.1.8 (Links). We introduce in Figure 2.1 three sorts of *links* $\text{---}\bullet$, $\otimes^n$ and $\mathfrak{M}^n$ — which correspond to MLL_u rules.

Every link may have two kinds of *ports*: *principal* ones, indexed by 0 and written below, and *auxiliary* ones, indexed by $1, \dots, n$ and written above. The auxiliary ports are ordered, but as we always represent the links as in Figure 2.1, we may safely omit the numbering. Axiom links have two principal ports, both indexed with 0, but we can always differentiate them (for instance by naming them 0_r and 0_l) if we need to. Remark that there is no cut link: a cut is represented with an edge between two principal ports.

Definition 2.1.9 (Proof nets for MLL_u). Given a proof of MLL_u , we may as usual build a *proof net* from it.

The *type* of a proof net $\mathcal{P}$ is Γ if there exists a proof of $\vdash \Gamma$ that infers it. A proof net always has several types, but up to α -equivalence (renaming of the literals) we may always assume it has a unique *principal type*. If a proof net may be typed with Γ , then for every A it may be typed with $\Gamma[A]$. By extension we will use the notion of type of an edge.

Remark. The same proof net—as it abstracts proofs— may be inferred by several proofs. Conversely, several graphs—as representations of proof nets— may correspond to the same proof net: we get round of this difficulty by associating to every proof net a single drawing among the possible drawings with the minimal number of crossings between edges, for the sake of simplicity. Two graphs representing proof nets that can be obtained from the same proof are taken to be equal.

TERUI uses a notion of *decorated proofs* to build his proof nets directly from proofs. It allows not to take into account *pseudo nets* and to avoid having to define a correction criterion. But it has a major drawback: in fact TERUI does not use proof nets to compute, but stick on this level of “description”. They are a rather convenient linear description of the objects we manipulate, but our setting allows us to skip this intermediate step, and to truly compute with proof nets.

Definition 2.1.10 (Size and depth of a proof net). The size $|\mathcal{P}|$ of a proof net $\mathcal{P}$ is the number of its links. The depth of a proof net is defined with respect to its type:

- The depth of a formula is defined by recurrence:

$$d(\alpha) = d(\alpha^\perp) = 1 \qquad d(\otimes^n(\vec{A})) = d(\mathfrak{M}^n(\overleftarrow{A})) = 1 + \max(d(A_1), \dots, d(A_n))$$

- The depth $d(\pi)$ of a proof π is the maximum depth of cut formulae in it.
- The depth $d(\mathcal{P})$ of a proof net $\mathcal{P}$ is $\min\{d(\pi) \mid \pi \text{ infers } \mathcal{P}\}$. The depth $d(\mathcal{P})$ of a proof net depends on its type, but it is minimal when we consider its principal type.

Now that we defined measurements on our objects, we can justify the use of *unbounded* Multiplicative Linear Logic. TERUI justifies this choice by writing that “MLL_u just gives us a depth-efficient way of writing proofs” and refer to the classical work on generalized multiplicatives previously mentioned [38]. This last approach is a classical work in Linear Logic and a major outbreak in the study of Linear Logic, but it focused on logical, mathematical and somehow combinatorial topics. The following remark explains from a complexity point of view this choice.

Remark (Unfolding MLL_u, folding MLL). We take the more “depth-efficient” way of encoding MLL into MLL_u, namely:

$$\otimes^i(A_1, \dots, A_i) \equiv (\dots((A_1 \otimes A_2) \otimes (A_3 \otimes A_4)) \otimes \dots \otimes ((A_{i-3} \otimes A_{i-2}) \otimes (A_{i-1} \otimes A_i)) \dots)$$

But we adopt MLL_u not only for convenience, for going back and forth between MLL_u and MLL has a cost. To unfold a proof net of MLL_u, we have to perform the two operations in Figure 2.2 for every $\wp^i$ - and $\otimes^i$ -link with $i > 2$. They are strictly local and thus can be performed with AC⁰ resources.

Doing the opposite operation —folding a MLL proof net into a MLL_u proof net— is a more global operation, for it has to identify the maximal chains of $\wp$ or $\otimes$ that can be folded and hence cannot be performed efficiently in parallel. But a really simple algorithm could be written to perform such folding, and it can be proved that log-s pace resources would be enough to execute it.⁵

It could be easily proved that this (un)folding operation is sound with respect to the type and the normalisation of the proof nets.

Concerning the bounds of such proof nets, let $\mathcal{P}$ be a MLL_u proof net, $d(\mathcal{P})$ its depth and $|\mathcal{P}|$ its size. Let $|c_i|_{\mathcal{P}}$ be the number of links with i auxiliary ports in $\mathcal{P}$, we have $|\mathcal{P}| \geq \sum_{i \geq 2} |c_i|_{\mathcal{P}}$. Now let $\mathcal{Q}$ be the MLL un-folded version of $\mathcal{P}$, it is clear that $|\mathcal{Q}| \geq \sum_{i \geq 2} |c_i|_{\mathcal{P}} \times \log(i)$. In the worst case the maximal arity of a link in $\mathcal{P}$ is $\mathcal{O}(|\mathcal{P}|)$, and we know that in this case $|\mathcal{Q}|$ is $\mathcal{O}(|\mathcal{P}| \times \log(|\mathcal{P}|))$. For the same reason, the depth of $\mathcal{Q}$ is in the worst case $\mathcal{O}(d(\mathcal{P}) \times \log(|\mathcal{P}|))$.

Roughly, we lose a lot regarding the depth if we stick to MLL proof nets, because we will take our proof nets to be of depth logarithmic in the size of the input. But as soon as the proof nets for MLL develop a computational power superior or equal to **L**, we may “for free” rewrite them as MLL_u proof nets.

Definition 2.1.11 (Cuts and parallel cut-elimination). A cut is an edge between the principal ports of two links. If one of these links is an *ax*-link, two cases occur:

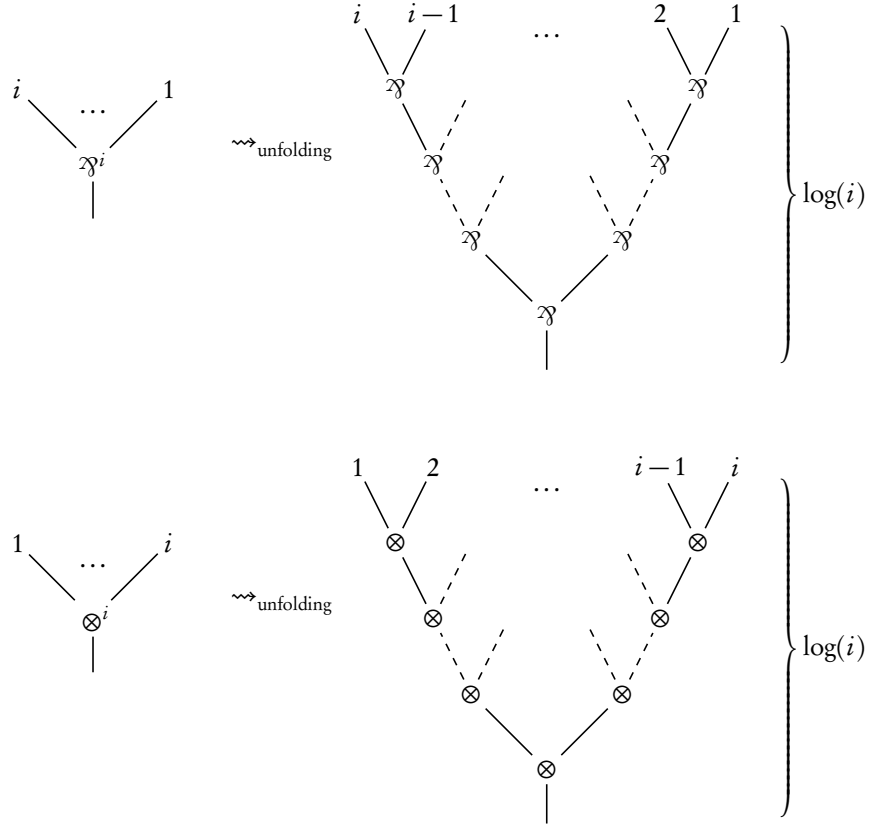
if the other link is an *ax*-link, we take the maximal chain of *ax*-links connected by their principal ports and defines this set of cuts as a *t-cut*,

otherwise the cut is an *a-cut*.

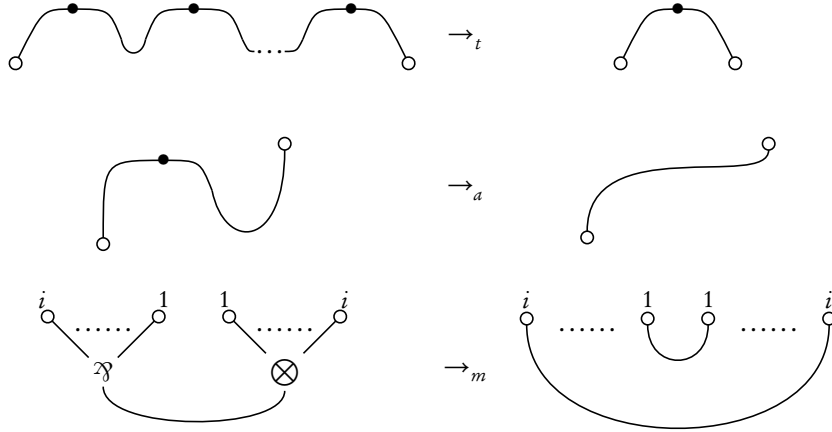
Otherwise it is a *m-cut* and we know that for $n \geq 2$, one link is a $\otimes^n$ -link and the other is a $\wp^n$ -link.

We define on Figure 2.3 three rewriting rules on the proof nets. For $r \in \{t, a, m\}$, if $\mathcal{Q}$ may be obtained from $\mathcal{P}$ by erasing all the r -cuts of $\mathcal{P}$ in parallel, we write $\mathcal{P} \rightrightarrows_r \mathcal{Q}$. If $\mathcal{P} \rightrightarrows_t \mathcal{Q}$, $\mathcal{P} \rightrightarrows_a \mathcal{Q}$ or

⁵Those two algorithms and complexity bounds would require to be stated formally the Direct Connection Language for proof nets, whose definition is to come. As it is not of major interest, we take the liberty to skip this red tape.

Figure 2.2: Unfolding a MLL_u proof net

For all $\circ \in \{(\wp^n)_{n \geq 2}, (\otimes^n)_{n \geq 2}\}$, $\circ$ may be $\bullet$ in $\rightarrow_m$.

Figure 2.3: Cut-elimination in parallel: t -, a - and m -reductions

$\mathcal{P} \Rightarrow_m \mathcal{Q}$, we write $\mathcal{P} \Rightarrow \mathcal{Q}$. To *normalise a proof net* $\mathcal{P}$ is to apply $\Rightarrow$ until we reach a cut-free proof net. We define $\Rightarrow^*$ as the transitive reflexive closure of $\Rightarrow$.

The t -cut is usually not defined, as it is just a special case of the a -cut. It was introduced by TERUI in order to be able to eliminate the cuts of a proof net in parallel. We know since “Linear logic” [47] that the number of steps needed to normalise a proof net is linear in its size. This is unsatisfactory for us in terms of parallelization: we would like the number of steps to be relative to the depth of the proof net, as it is the case for the Boolean circuits.

If we try roughly to reduce in parallel a cut between two ax -links, we are faced with a critical pair, as exposed in Figure 2.4. TERUI avoids this situation by using this t -reduction (*tightening reduction*) which eliminates in one step all those critical pairs. Of course, this rule re-introduce some sequentialization in the parallel cut-elimination⁶ but after applying it, we can then safely reduce all the other cuts in parallel.

Theorem 2.1.1 (Parallel cut-elimination TERUI [130]). *Every proof net $\mathcal{P}$ normalises in at most $\mathcal{O}(d(\mathcal{P}))$ applications of $\Rightarrow$.*

The idea of the proof is quite simple: after the application of $\Rightarrow_t$ and $\Rightarrow_a$, no t -cut or a -cut remains in the proof net, because $\Rightarrow_a$ cannot create a t -cut. So either the proof net obtained is cut-free, and the proof is done, either some m -cuts are present. In this latter case, an application of $\Rightarrow_m$ eliminates at least one m -cut and the depth of the whole proof net lowers by 1.

By iterating the procedure, we get the expected result. This procedure allows us to make the most of the computational power of proof nets, by achieving a “speed-up” in the number of steps needed to normalise them.

Now that we know that one could compute efficiently, in parallel, with proof nets, we should wonder *what* we can compute!

2.2 Boolean Proof Nets

We will begin with a reminder of *how* a boolean value may be represented and computed with proof nets. Next we will introduce our contribution, proof circuits, and make sure that they are proof nets.

Proof Nets that compute Boolean functions

We just saw that proof nets are accompanied by a mechanism of evaluation, but this is not sufficient to represent actual computation. We have now to define how proof nets represent Boolean values (Definition 2.2.1) and Boolean functions (Definition 2.2.2). To study them in a uniform framework we define their Direct Connection Language (Definition 2.2.3). This will allow us later on to use them as “true” model of computation, i.e. a “reasonable” model if we remember the *Invariance Thesis*.

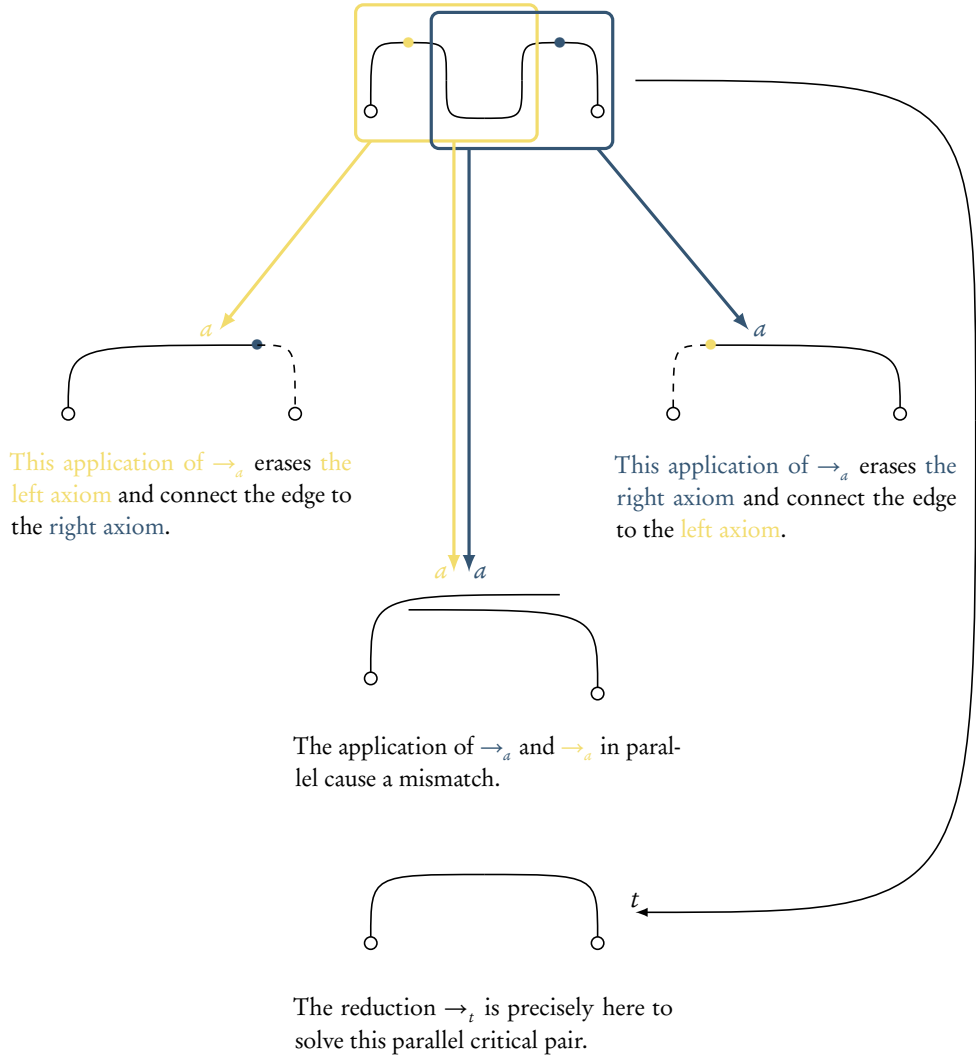
We cannot in MLL_u derive the classical representation of the Boolean values, i.e. $\alpha \multimap \alpha \multimap \alpha$, but we use its ternary variant, without weakening.

Definition 2.2.1 (Boolean type, b_0 and b_1 TERUI [130]). Let $\mathbf{B}$ be the boolean type of MLL_u :

$$\mathbf{B} := \wp^3(\alpha^\perp, \alpha^\perp, (\alpha \otimes \alpha))$$

Its proof is

⁶The process of identifying maximal chains of axioms with cuts between them is necessarily done sequentially.



The proof net at the top of the figure has a single cut between two axioms, but it can be reduced in two different ways, depending on the axiom we “keep”. This is illustrated by the two rectangles, which emphasize the two premisses of the application of $\rightarrow_a$. If we perform the two possible applications in parallel, we do not get the expected result.

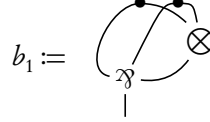
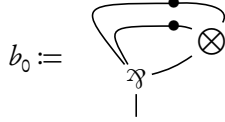
The t -r eduction will handle this situation, but it requires the USTConn_2 function to be performed in parallel, in order to identify the maximal chains of axioms with a cut between them.

Figure 2.4: An obstacle to cut-elimination in parallel

$$\frac{\frac{\overline{\vdash \alpha, \alpha^\perp} \text{ ax.} \quad \overline{\vdash \alpha, \alpha^\perp} \text{ ax.}}{\vdash \alpha^\perp, \alpha^\perp, (\alpha \otimes \alpha)} \otimes^2}{\vdash \mathfrak{X}^3(\alpha^\perp, \alpha^\perp, (\alpha \otimes \alpha))} \mathfrak{X}^3$$

But if we look closely, this proof is not unique: does the first literal $\alpha^\perp$ of the $\mathfrak{X}$ connective comes from the right hand-side axiom or from the left hand-side axiom⁷?

Depending on the answer to this question, we obtain the two proof nets of type **B**: b_0 and b_1 , respectively used to represent false and true.



We write $\vec{b}$ for $b_{i_1}, \dots, b_{i_n}$ with $i_1, \dots, i_n \in \{0, 1\}$.

As we can see, b_0 and b_1 differ only on their planarity: the proof nets exhibit the exchange that was kept implicit in the proof. This difference of planarity is the key of the encoding of the boolean values into proof nets.

We present briefly two examples of computation: [Figure 2.5](#) shows how to inverse the planarity and [Figure 2.6](#) how conditional can be simulated thanks to planarity.

Definition 2.2.2 (Boolean proof nets). A *Boolean proof net with n inputs* is a proof net $P_n(\vec{p})$ of type

$$\vdash \mathbf{B}^\perp[A_1], \dots, \mathbf{B}^\perp[A_n], \otimes^{1+m}(\mathbf{B}[A], D_1, \dots, D_m)$$

The tensor in this type is the *result tensor*: it collects the result of the computation on its first auxiliary port and the *garbage* —here of type $D_1, \dots, D_m$ — on its other auxiliary ports which will be called the *garbage ports*.

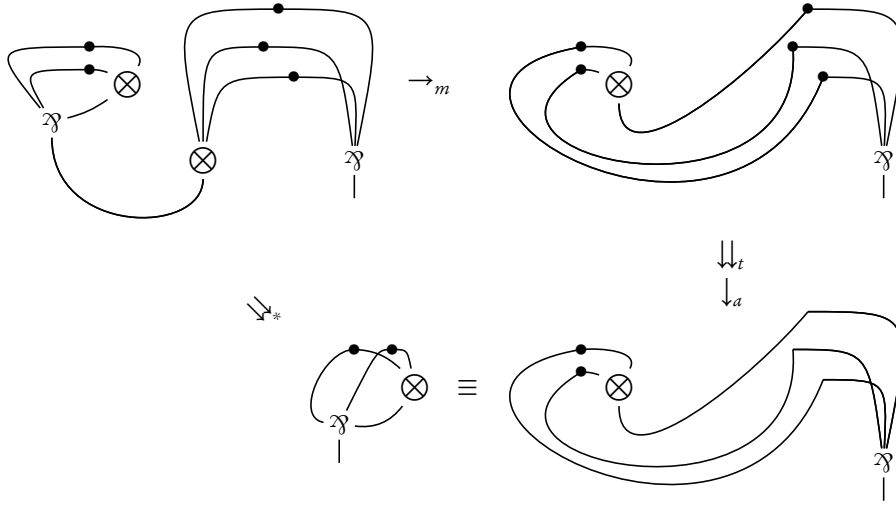
Given $\vec{b}$ of length n , $P_n(\vec{b})$ is obtained by connecting with cuts $\mathbf{B}^\perp[A_j]$ to b_{i_j} for all $1 \leq j \leq n$. We have that $P_n(\vec{b}) \Rightarrow^* Q$ where Q is unique, cut-free and of type $\otimes^{1+m}(\mathbf{B}[A], D_1, \dots, D_m)$. Let b_r be the proof net of type **B** connected to the first port of $\otimes^{1+m}$. We know since [Definition 2.2.1](#) that the only two⁸ normal proof nets of type **B** are b_0 and b_1 , so $r \in \{0, 1\}$ is considered as the result of this computation, and we write $P_n(\vec{b}) \rightarrow_{\text{ev.}} b_r$, see [Figure 2.7](#) for a better insight.

As for the Boolean circuits, a Boolean proof net can only receive a fixed size of input, so we need to work with families of Boolean proof nets $\mathbf{P} = (P_n)_{n \in \mathbb{N}}$. We say that a Boolean proof net family $\mathbf{P}$ represents a Boolean function f^n if for all $w \equiv i_1 \dots i_n \in \{0, 1\}^n$,

$$P_n(b_{i_1}, \dots, b_{i_n}) \rightarrow_{\text{ev.}} b_{f(w)}$$

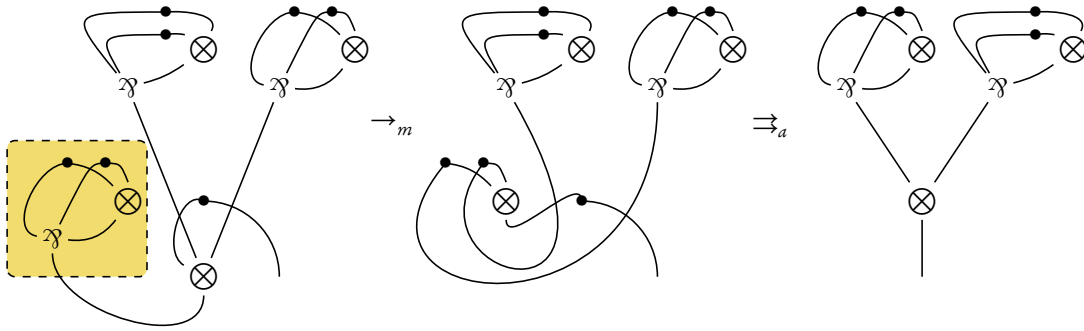
⁷This ambiguity rests of course on the implicit exchange. In fact, there could be more exchanges, but we can up to reordering consider only this one.

⁸Up to η -expansion, but we can safely ignore this fact by assuming that the type chosen is minimal, i.e. that b_r is of type **B** and not of type $\mathbf{B}[C]$ for any $C \neq \alpha$.



The proof net we present is the result of the “branching”, i.e. the connection *via* a cut, of b_0 with a proof net that will be defined later on as $\mathcal{N}_{\text{eg}}$. This proof net adds during the cut-elimination a cross between two edges. One can simply check that if the input was b_1 , this proof net would have removed the crossing between the edges. Therefore, this small proof net is the function mapping 0 to 1 and 1 to 0: the negation.

Figure 2.5: Adding or removing a crossing: how negation works



The input (here in a dashed rectangle) is a proof net of type $\mathbf{B}[\mathbf{B}]$, and it “selects” —according to its planarity or non-planarity— during the normalisation which one of b_0 or b_1 is connected to the first auxiliary port of the tensor and will be considered as the result —the other being treated as garbage.

We will in the following replace b_0 and b_1 with more complex proof nets, let’s say $\mathcal{P}$ and $\mathcal{Q}$, and we can remark that this proof net compute “if the input is true then $\mathcal{Q}$ else $\mathcal{P}$ ”.

Figure 2.6: The conditional, the core of the computation

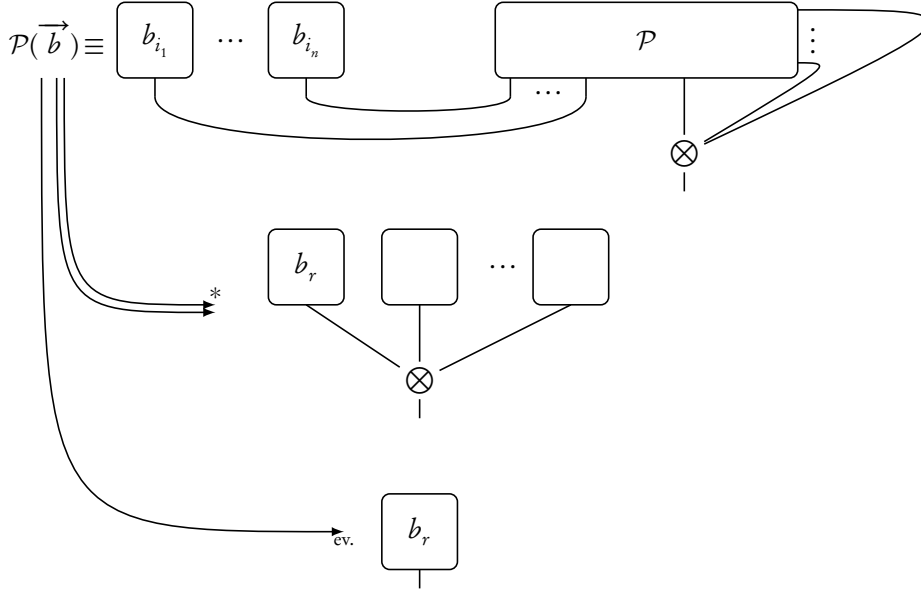


Figure 2.7: Branching variables on a Boolean proof net and evaluating it

Remark that the Boolean type may contains more than literals, thanks to the substitution. Without substitutions, Boolean proof nets would be *flat*, i.e. necessarily of type

$$\mathbf{B}^\perp[\alpha], \dots, \mathbf{B}^\perp[\alpha], \otimes^{1+m}(\mathbf{B}[\alpha], D_1, \dots, D_m)$$

and TERUI proved that in this case we could not compute more than *Parity*ⁿ or its negation!

We may easily define *language accepted by a family of Boolean proof nets* and *uniformity*, thanks to the definition of Direct Connection Language for proof nets that follows.

Definition 2.2.3 (Direct Connection Language for proof nets MOGBIL and RAHLI [107]). Given $\mathbf{P} = (P_n)_{n \in \mathbb{N}}$ a family of Boolean proof nets, its *Direct Connection Language* $L_{DC}(\mathbf{P})$ is the set of tuples $\langle 1^y, \bar{g}, \bar{p}, \bar{b} \rangle$ where g is a link in P_y , of sort b if $p = \epsilon$ else the p th premise of g is the link b .

If $\langle 1^y, \bar{p}, 0, \bar{b} \rangle$ or $\langle 1^y, \bar{b}, 0, \bar{p} \rangle$ belong to $L_{DC}(\mathbf{P})$, there is a cut between b and p in P_y .

Of course this definition is really close to the definition of Direct Connection Language for Boolean circuits (Definition 2.1.3), for the object are in fact very similar. We will make them even more similar by constraining how Boolean proof nets are built.

Proof Circuits

Proof circuits were introduced in my *mémoire de Master* [7] and appears in a paper published in the Proceedings of the Second Workshop on Developments in Implicit Computational Complexity [8]. They make it simpler to grasp the mechanisms of computation at work in the Boolean proof nets. They are studied in a uniform framework, following the work initiated by MOGBIL and RAHLI [107, 105], but they help to avoid a complex and tedious composition and are closer to Boolean circuits; also strictly speaking, for we extend the correspondence between them to sublogarithmic classes of complexity.

A proof circuit is a Boolean proof net (Theorem 2.2.3) made out of pieces (Definition 2.2.4) which represents Boolean functions, constants or duplicates values. Garbage is manipulated in an innovative way, as shown in Figure 2.9. The mechanisms of computation explained in Figure 2.5 and Figure 2.6 remain valid in this setting. Apart from the constants and the $\mathcal{N}_{\text{eg}}$ piece, all pieces use the conditional presented in Figure 2.6 with a “built-in” composition.

From now on every edge represented by  is connected on its right to a free (unconnected) garbage port of the result tensor: it carries a piece of garbage. As there is no modality in MLL_{u} , the computation cannot erase the useless values generated during the course of the computation, so we handle them differently.

Patrick BAILLOT argued that proof circuits should rather be seen as a target to the translation from Boolean circuits to Boolean proof nets, as a technical tool of a proof, rather than as an object to define classes of complexity. It is true that the original motivation that led to introduce them was to sharpen the translation to be developed in Section 2.3. Yet, as a by-product, they allow to define a framework where the fan-in of the *pieces* is bounded, and so to get an equivalent to the bounded Boolean circuits. It is true that their nature makes them look rather as Boolean circuits, but the bridging they propose between those two worlds is the key to the new correspondence developed in Section 2.4.

Definition 2.2.4 (Pieces). We present in Table 2.8 the set of *pieces* at our disposal. Entries are labelled with e , exits with s and garbage with g . Edges labelled b_k —for $k \in \{0, 1\}$ — are connected to the edge labelled s of the piece b_k .

A piece $\mathcal{P}$ with $i \geq 0$ entries, $j \geq 1$ exits and $k \geq 0$ garbage is one of the pseudo nets in Table 2.8, where i edges are labelled with $e_1, \dots, e_i$, j edges are labelled with $s_1, \dots, s_j$ and k edges labelled with $g_1, \dots, g_k$ go to the garbage ports of the result tensor.

Table 2.8: Pieces of Proof Circuits

We set $2 \leq j \leq i$. If $i = 2$, the edge a is the edge s in the pieces $\mathcal{D}_{\text{isj}}^i$ and $\mathcal{C}_{\text{onj}}^i$.

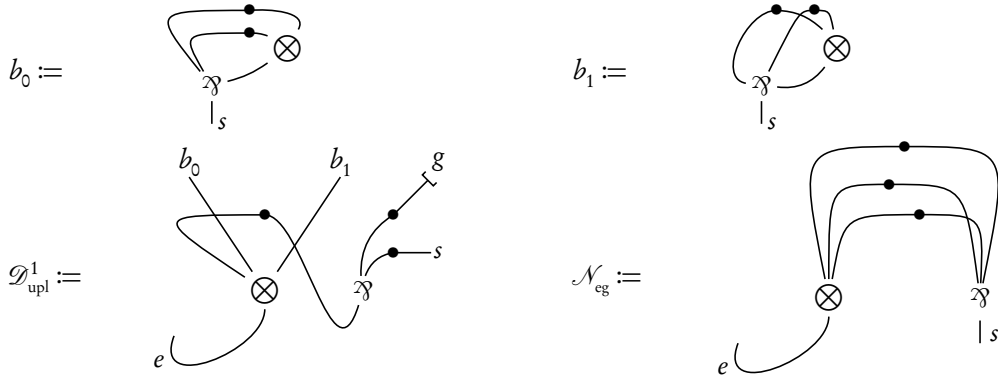
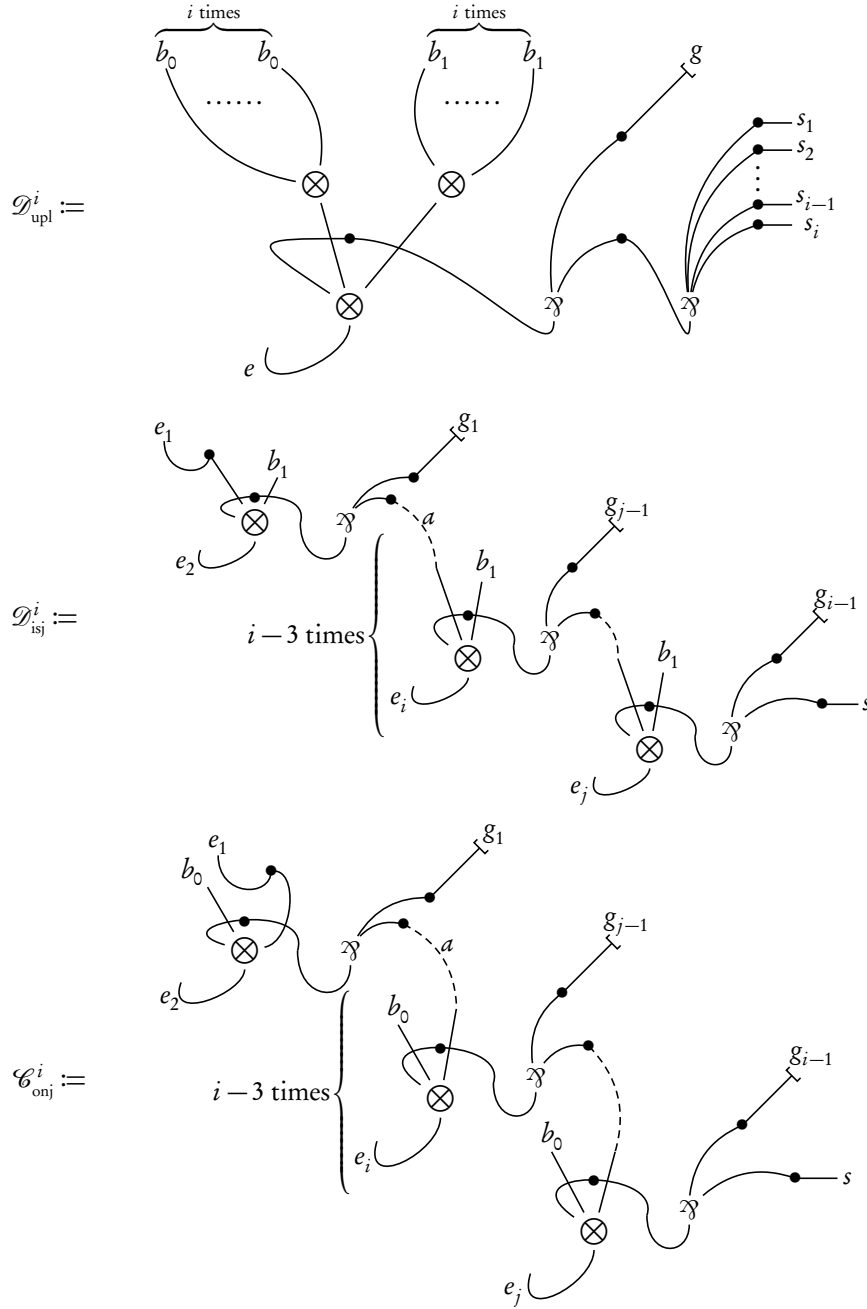


Table 2.8: Pieces – continued from previous page



We define two sets of pieces:

$$\mathcal{P}_u = \{b_0, b_1, \mathcal{N}_{\text{eg}}, \{\mathcal{D}_{\text{upl}}^i\}_{i \geq 1}, \{\mathcal{D}_{\text{isj}}^i\}_{i \geq 2}, \{\mathcal{C}_{\text{onj}}^i\}_{i \geq 2}\}$$

and

$$\mathcal{P}_b = \{b_0, b_1, \mathcal{N}_{\text{eg}}, \{\mathcal{D}_{\text{upl}}^i\}_{i \geq 1}, \mathcal{D}_{\text{isj}}^2, \mathcal{C}_{\text{onj}}^2\}$$

To compose two pieces $\mathcal{P}_1$ and $\mathcal{P}_2$, we connect an exit of $\mathcal{P}_1$ to an entry of $\mathcal{P}_2$. It is not allowed to loop: we cannot connect an entry and an exit belonging to the same piece. All the same, we cannot have $\mathcal{P}_1$ connected to $\mathcal{P}_2$ connected to $\mathcal{P}_3 \dots$ connected to $\mathcal{P}_1$.

An entry (resp. an exit) that is not connected to an exit (resp. an entry) of another piece is said to be *unconnected*.

A piece with i entries and j exits will sometimes be referred to as having *fan-in* i and *fan-out* j , as for the gates of the Boolean circuits.

Definition 2.2.5 (Proof circuits). A proof circuit $P_n(\vec{p})$ with n inputs and one output is obtained by composing pieces such that n entries and one exit are unconnected. If no garbage is created we add a $\mathcal{D}_{\text{upl}}^1$ piece connected to the unconnected exit to produce some artificially. Then we add a result tensor whose first auxiliary port is connected to the exit—which is also the output of the proof circuit—and whose garbage ports are connected to the garbage of the pieces. We then label every unconnected entries with $p_1, \dots, p_n$: those are the inputs of the proof circuit.

Given $\vec{b}$ of length n , $P_n(\vec{b})$ is obtained by connecting with cuts p_j to b_{i_j} for all $1 \leq j \leq n$. The Figure 2.7 can be used as it to describe the branching of the inputs.

Proof circuits come in two flavours: they may be composed of *unbounded* pieces (taken in the set $\mathcal{P}_u$ or of *bounded* (or *binaries*) pieces (taken from the set $\mathcal{P}_b$), as for the Boolean circuits over the basis $\mathcal{B}_u$ or $\mathcal{B}_b$. Yet the Linear Logic we use is in both cases the unbounded MLL_u, and the $\mathcal{D}_{\text{upl}}^i$ pieces can be of any fan-out.

Without loss of generality, one may consider that the $\mathcal{N}_{\text{eg}}$ piece are only connected to the inputs of the proof circuit and that a $\mathcal{D}_{\text{upl}}^i$ piece is never connected to another $\mathcal{D}_{\text{upl}}^i$ piece: we do not lose computational power if we negate only the entries, and we don't need to duplicate i times the values that have just been duplicated j times, for we might as well duplicate it $i + j$ times directly. This comes from a semantic of evaluation of those proof circuits that will come soon.

There are two main novelties in this way to define Boolean proof nets:

First, the composition is handled by the pieces themselves. If you recall Figure 2.6, you may notice that the value computed, the output, is connected to this first port of the tensor. So to access this value, one has to cut the tensor with a par, to “extract” the output to be able to compute with it. Here the garbage is evacuated as soon as it is produced. We present in Figure 2.9 the difference of representation of the disjunction of arity 4 between TERUI's framework and our setting.

We define directly proof nets that can be built easily and where mechanisms of computation are simpler to grasp. This allow us to lower the computational power needed to build them, as we will see in Section 2.3.

But first we must make sure that they are proof nets, and for that we will need the following definition.

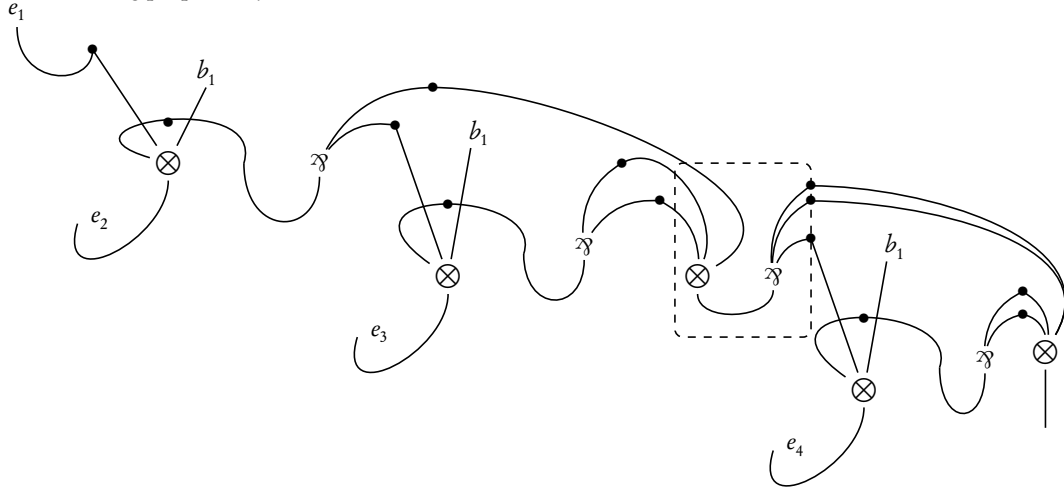
Definition 2.2.6 (Height of a piece). Let P be a proof circuit, we define for any of its pieces $\mathcal{P}$, its *height* denoted by $h(\mathcal{P})$:

If the exit of $\mathcal{P}$ is connected to the result tensor, then $h(\mathcal{P}) = 0$.

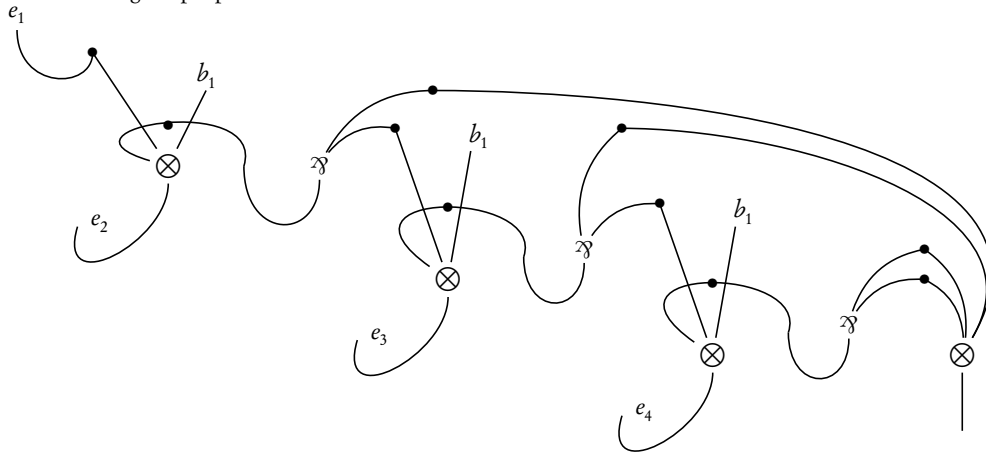
If $\mathcal{P}$ has a single exit and it is connected to an entry of a piece $\mathcal{P}'$, then $h(\mathcal{P}) = h(\mathcal{P}') + 1$.

We give below two encodings of $\vee^4$, the disjunction of arity 4.

The encoding proposed by TERUI:

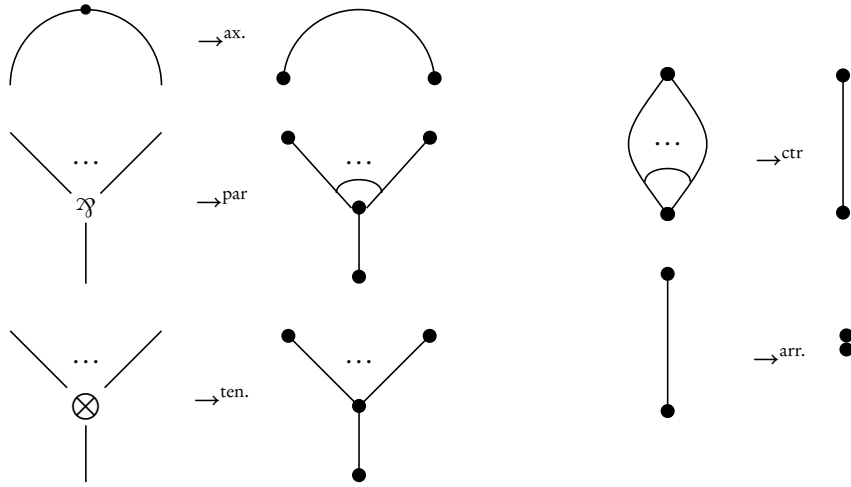


The encoding we propose:



Note that the couple $\otimes/\wp$ in the first encoding (in a dashed rectangle) disappears in the second encoding, exactly as if we had reduced this cut. It might look quite a minimal improvement, but it has to be conceived at the scale of the whole Boolean proof net. We will show in the proof of [Theorem 2.3.3](#) that it allows us to drastically lower the size of the manipulated object.

Figure 2.9: The improvement of “built-in” composition



The three rules on the left goes from pseudo nets to dotted graphs, the two on the right from dotted graphs to dotted graphs. The dots must be different in the case of $\rightarrow^{\text{ctr}}$ and $\rightarrow^{\text{arr.}}$. In the $\rightarrow^{\text{arr.}}$ case, there should be no semi-arc coming from the relabelling of a $\mathcal{P}$.

Figure 2.10: Contractibility rules: relabelling and contracting a pseudo net

If $\mathcal{P}$ has $i > 1$ exits connected to entries of $\mathcal{P}_1, \dots, \mathcal{P}_i$, then $h(\mathcal{P}) = \max_{n \leq i} (h(\mathcal{P}_n)) + 1$.

Why this definition ? For two reasons: as we don't know yet that proof circuits are proof nets, we do not know that they may be typed, and yet we need to be able to measure the degree of “overlapping” of their pieces. The second reason is that the depth do not increase each time we cross a piece, as it will be shown in [Lemma 2.3.1](#), but we need a criterion to reason inductively with. We use this notion in the following subsection, to prove [Theorem 2.2.3](#).

Contractibility criterion: Proof Circuits are Boolean Proof Nets

For the time being, proof circuits are only freely-generated by connecting pieces, and we have no insurance whatsoever that they are actually proof nets. They are only pseudo nets, for we don't know yet if they are typable, if every one of them may be linearized as a proof.

To prove that all proof circuit is a proof net, we use a generalized version of the contractibility criteria of [DANOS \[36\]](#). We adopt its variant introduced by [LAFONT \[93\]](#) and refined by [GUERRINI and MASINI \[67\]](#) who use *parsing boxes*.

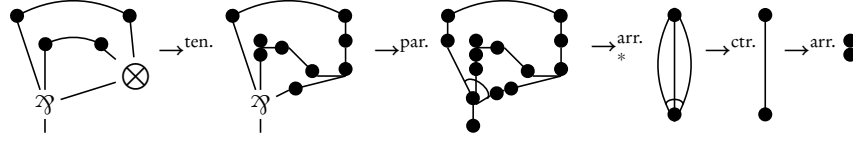
The contraction rules —adapted to the unbounded case of MLL_u — are reminded in [Figure 2.10](#). Those rules have a different role: some of them ($\rightarrow^{\text{ax.}}$, $\rightarrow^{\text{par}}$ and $\rightarrow^{\text{ten.}}$) transform a pseudo net into a dotted graph —i.e. they are only a relabelling of the pseudo net—, and the others ($\rightarrow^{\text{arr.}}$, $\rightarrow^{\text{ctr.}}$) contract the dotted graph, constitutes a rewriting system. Let $\rightarrow^r$ be one of those rules, we write $\rightarrow^r_*$ its transitive reflexive closure.

There is no need to define a relabelling for the cuts, for links eventually transforms into dots, and an edge between two dots —that used to be a cut— will be contracted with a single application of $\rightarrow^{\text{arr.}}$.

Theorem 2.2.1 (Contractibility criterion, adapted from DANOS [36]). *A pseudo net is a proof net of MLL_u iff after applying $\rightarrow^{ax.}$, $\rightarrow^{par.}$ and $\rightarrow^{ten.}$ as many time as possible, the dotted graph obtained can be reduced to a single dot applying $\rightarrow^{arr.}$ and $\rightarrow^{ctr.}$.*

We will in the following write that a pseudo net “reduces” or “contracts” to a dot to express that it is possible to obtain a dot from it by applying the rules of Figure 2.10.

Definition 2.2.7 ($\rightarrow_*^b$). The proof nets representing the boolean values, b_0 and b_1 , reduce themselves easily into a dot. After an application of $\rightarrow_*^{ax.}$, b_0 reduces itself into a dot by applying the following sequence:



The case for b_1 is the same up to a minor permutation of two wires, and we define $\rightarrow_*^b$ as the reduction in one step of all the pieces b_0 and b_1 into dots.

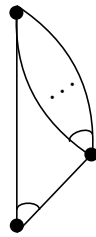
Lemma 2.2.2. *A “piece” whose exits and garbage are connected to a single dot⁹ reduces to a dot.*

Proof. Let $\mathcal{P}^*$ be the piece $\mathcal{P}$ where all exits and garbage are connected to the same dot. We study the different cases:

b_0 and b_1 After applying $\rightarrow^b$ and $\rightarrow^{arr.}$, b_0^* as b_1^* contracts to a dot.

$\mathcal{N}_{eg}^*$ After applying $\rightarrow_*^{ax.}$, $\rightarrow^{ten.}$, $\rightarrow^{par.}$, $\rightarrow^{arr.}$, $\rightarrow^{ctr.}$ and $\rightarrow^{arr.}$, $\mathcal{N}_{eg}^*$ reduces to a dot.

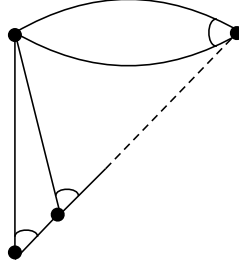
$\mathcal{D}_{upl}^i$ After applying $\rightarrow_*^b$, $\rightarrow_*^{ax.}$, $\rightarrow_*^{ten.}$, $\rightarrow_*^{arr.}$, $\rightarrow_*^{par.}$ and then $\rightarrow_*^{arr.}$, $\mathcal{D}_{upl}^{i*}$ is reduced to



If we apply $\rightarrow^{ctr.}$, $\rightarrow^{arr.}$, $\rightarrow^{ctr.}$ and $\rightarrow^{arr.}$, we get the expected result.

$\mathcal{C}_{onj}^i$ and $\mathcal{D}_{isj}^i$ After applying $\rightarrow_*^b$, $\rightarrow_*^{ax.}$, $\rightarrow_*^{ten.}$, $\rightarrow_*^{par.}$ and $\rightarrow_*^{arr.}$, $\mathcal{C}_{onj}^{i*}$ as well as $\mathcal{D}_{isj}^{i*}$ is

⁹Note that we are speaking of a structure midway between the pseudo net and the dotted graph. It is not strictly speaking a piece, but this lemma is suited to prove Theorem 2.2.3 later on.



If we apply $i - 1$ times the pattern $\rightarrow^{\text{ctr.}}, \rightarrow^{\text{arr.}}$, we get the expected result.

Q.E.D.

Theorem 2.2.3. *Every proof circuit is a Boolean proof net.*

Proof. We will prove this by using the height of the pieces (recall Definition 2.2.6). Given P a proof circuit, every of his piece has a height $n \in \mathbb{N}$, we prove that the set of pieces of height equal or inferior to n contracts itself into a dot.

If $n = 0$ The set of pieces whose height is 0 is by definition a singleton. His exit and his garbage are connected to the result tensor. There might be other garbage that come from piece of height > 0 connected to the result tensor, but we may safely ignore them. All we want is to fulfill the hypothesis of Lemma 2.2.2, and this is the case after applying $\rightarrow^{\text{ten}}$ and $\rightarrow_{*}^{\text{arr.}}$, so this piece contracts itself into a dot.

If $n > 0$ By definition, all pieces of height n are connected to pieces of height strictly inferior to n . By induction, we know that this set of pieces of height inferior to n contracts itself into a single dot. So all pieces of height n have their exits and garbage connected to a single dot, and by Lemma 2.2.2 this set contract itself into a dot.

A proof circuit is made of pieces and of a result tensor, there always exists a n such that all its pieces are of height inferior or equal to n , so it contracts itself into a single dot. The Theorem 2.2.1 gives us that all proof circuits are proof nets.

Regarding the type of the proof circuit P : we know that there exists several derivations that infer it. The only conclusions are:

- unconnected entries of pieces, that is input,
- the result tensor.

We know that the type of an entry of a piece is —up to substitution— $\mathbf{B}^{\perp}$ and that the result tensor recollect on its first port a formula of type $\mathbf{B}$. Hence, any proof circuit can be typed with

$$\vdash \mathbf{B}^{\perp}[A_1], \dots, \mathbf{B}^{\perp}[A_n], \otimes^{1+m}(\mathbf{B}[A], D_1, \dots, D_m)$$

so it is a Boolean proof net.

Q.E.D.

Patrick BAILLOT questioned this use of the contractibility criterion. His argument is that as proof circuits are defined inductively, one could directly build its derivation by induction. This argument is perfectly valid, and exhibiting concretely the derivation of a proof circuit would give us almost for free [Corollary 2.3.2](#) to come. Yet, we maintain that this proof technique is the “right” one, because it highlights the graphical structure of our object, its parallel nature. Moreover, it dispenses us from maintaining a high stack of variables for the number of inputs, the fan-in and fan-out of the piece under study, the size of the result tensor, etc. Finally, it seems to us that this proof is in better adequation with the situation, for it testifies of the role of the result tensor to “recollect” all the pending conclusion regarding garbage.

Anyway, this theorem establishes that proof circuits are proof nets, so we can use the same evaluation mechanism, and benefit from [Theorem 2.1.1](#) regarding the parallel cut-elimination. We can define *families of proof circuits* $\mathbf{P} = (P_n)_{n \in \mathbb{N}}$ as for the Boolean circuits, and the notions of evaluation and representation of a Boolean function can be used “as it” from [Definition 2.2.2](#). As proof circuits are just a special case of Boolean proof nets, we can adopt for them the same Direct Connection Language we introduced in [Definition 2.2.3](#). At last, concerning uniformity, simply substitutes “gate” by “link” and “C” by “P” in [Definition 2.1.4](#) to obtain the suitable definition.

This theorem also mean that a type or a formula can be associated to every edge of a proof circuit. In the case of the entry or the exit of a piece, we will speak of a *transported value*: we know that this edge “carries” a boolean type that can be associated to b_0 or b_1 . This notion will turn out to be crucial when we will parse proof circuits with Alternating Turing Machines to determine what are the values carried along the edges without normalising the proof circuit. We refer to [Section 2.4](#) and more precisely to [Theorem 2.4.3](#) for a demonstration of the utility of this notion.

We are now ready to study the computational power of proof circuits: we quickly expose how to compose two proof circuits and then we define a *Proof Circuit Complexity* (PCC) class of complexity.

Remark. To compose two proof circuits P_1 and P_2 , we remove the result tensor of P_1 , identify the unconnected exit of P_1 with the selected input of P_2 , and recollect all the garbage with the result tensor of P_2 . We then label the unconnected entries anew and obtain a proof circuit. The result of this composition is illustrated in [Figure 2.11](#).

Definition 2.2.8 ($\mathbf{PCC}^i$ (resp. $\mathbf{mBN}^i$, MOGBIL and RAHLI [[107](#)])). A language $X \subseteq \{0, 1\}^*$ belongs to the class $\mathbf{PCC}^i$ (resp. $\mathbf{mBN}^i$) if X is accepted by a polynomial-size, $\log^i$ -depth uniform family of proof circuits (resp. of Boolean proof nets).

If the pieces are taken from $\mathcal{P}_b$, we write $\mathbf{bPCC}^i$ ¹⁰, elsewhere the fan-in of the pieces is unbounded.

Remark. For all $i \in \mathbb{N}$, $\mathbf{bPCC}^i \subseteq \mathbf{PCC}^i \subseteq \mathbf{mBN}^i$.

The first inclusion comes from the fact that every piece of $\mathcal{P}_b$ is a piece of $\mathcal{P}_u$. The second comes from [Theorem 2.2.3](#).

Remark that $\mathbf{PCC}^i \subseteq \mathbf{bPCC}^{i+1}$ with the same argument that proves $\mathbf{AC}^i \subseteq \mathbf{NC}^{i+1}$: every piece of fan-in $n > 2$ can be replaced by $n - 1$ pieces of fan-in 2, organized in a “depth-efficient way”,¹¹ thus multiplying the total depth of the proof circuit by a logarithm.¹²

We are now ready to compare proof circuits with other models of computation: first come Boolean circuits, and in the next section we will study the links with Alternating Turing Machines.

¹⁰See [errata](#).

¹¹With the same arrangement that the “unfolding procedure” of [Figure 2.2](#).

¹²Another encoding, proving that $\mathbf{PCC}^i = \mathbf{bPCC}^i$, is given in the [errata](#).

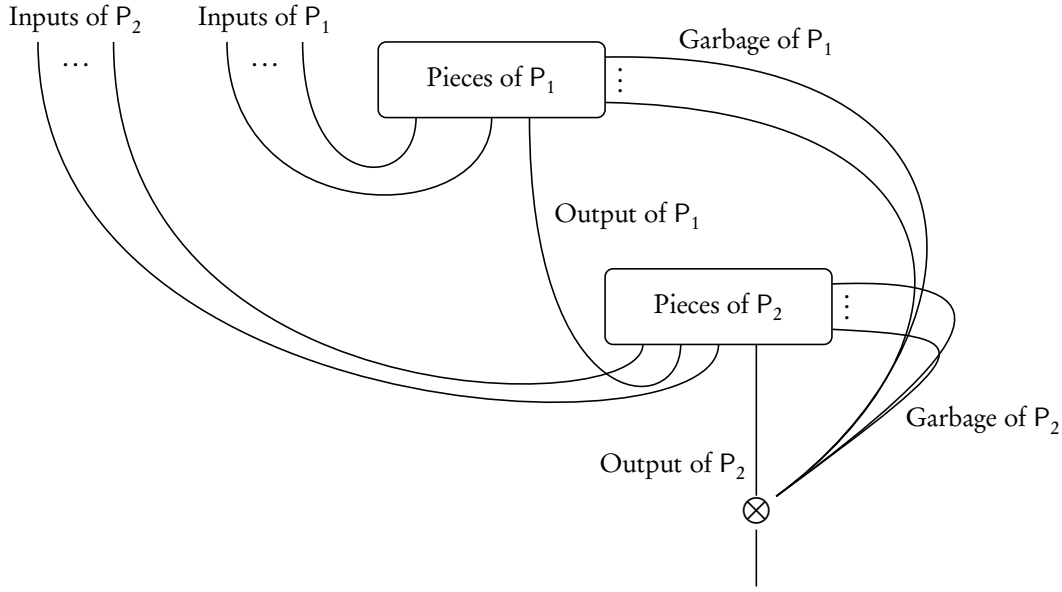


Figure 2.II: Composing two proof circuits

2.3 Correspondence With Boolean Circuits

The first part of this section is an improvement compared to the previous works about Boolean proof nets, the second is just a reminder of those previous works.

Translation from Boolean Circuits to Proof Circuits

By using our proof circuits we prove anew the translation from Boolean circuits to proof circuits, but lower its complexity and by doing so, we extend the results to sublogarithmic classes of complexity. We begin with an intermediate lemma and one of its corollaries.

Lemma 2.3.I. *Let P be a proof circuit, $\mathcal{P}$ one of its piece with $n \geq 1$ entries and $m \geq 1$ exits. Let $A_{e_1}, \dots, A_{e_n}$ (resp. $A_{s_1}, \dots, A_{s_m}$) be the type associated to its entries (resp. exits), we have for $j \in \{1, \dots, n\}$*

$$d(A_{e_j}) \leq d(\mathbf{B}) + \max(d(A_{s_1}), \dots, d(A_{s_m}))$$

Proof. We simply consider the different cases and refers to Table 2.8:

- $\mathcal{P} \notin \{b_0, b_1\}$ according to the hypothesis.
- If $\mathcal{P} = \mathcal{N}_{\text{eg}}$, $d(A_{e_1}) = d(A_{s_1})$.
- If $\mathcal{P} = \mathcal{D}_{\text{isj}}^i$ or $\mathcal{C}_{\text{onj}}^i$, then $d(A_{e_1}) = d(A_s)$ and $d(A_{e_j}) = d(\mathbf{B}) + d(A_s)$ for $2 \leq j \leq i$.

◦ If $\mathcal{P} = \mathcal{D}_{\text{upl}}^i$, then

$$\begin{aligned}
 d(A_{e_1}) &= d(\otimes^3((\overleftarrow{\mathfrak{A}}^n(A_{s_n}^\perp) \mathfrak{A}(\overleftarrow{\mathfrak{A}}^n(A_{s_n}^\perp)), \otimes^n(\overrightarrow{A_{s_n}}), \otimes^n(\overrightarrow{A_{s_n}}))) \\
 &= 1 + \max(d((\overleftarrow{\mathfrak{A}}^n(A_{s_n}^\perp) \mathfrak{A}(\overleftarrow{\mathfrak{A}}^n(A_{s_n}^\perp))), d(\otimes^n(\overrightarrow{A_{s_n}}))) \\
 &= 2 + \max(d(\otimes^n(\overrightarrow{A_{s_n}}))) \\
 &= 3 + \max(d(A_{s_1}), \dots, d(A_{s_n}))
 \end{aligned}$$

As $d(\mathbf{B}) = d(\mathfrak{A}^3(\alpha^\perp, \alpha^\perp, (\alpha \otimes \alpha))) = 3$, the lemma is proved.

Q.E.D.

The “less or equal” sign should not misguide the reader: this lemma teaches us that the depth slowly *increase* from the exits to the entries i.e. that the formulae associated to the entries are bigger than the formulae associated to the exits. We should also remark that the depth increases independently from the fan-in of the piece and that duplication has an impact on the depth. But since we cannot “duplicate more than we compute”, the impact will only be a multiplication by a constant factor (see proof of [Theorem 2.3.3](#) for a precise statement).

Corollary 2.3.2. *For all proof circuit $P_n(\overrightarrow{p})$ and all $\overrightarrow{b}$, the cuts at maximum depth in $P_n(\overrightarrow{b})$ are between the entry of a piece and an input b_{i_j} for some $1 \leq j \leq n$.*

Proof. The formulae in the garbage are never cut, and in all pieces the cuts that do not connect an entry or an exit (“the inner cuts”) are always of depth inferior or equal to cuts connecting the entries.

We just proved that the depths of the cut formulae slowly increase from the exit to the entry, and as the entries that are not connected to other pieces are connected to values, this lemma is proved.

Of course, an entry could also be connected to a constant b_0 or b_1 , but we may always consider that two constants b_0 and b_1 , eventually duplicated, are part of the input, as for the Boolean circuits.

Q.E.D.

Problem C: Translation from Boolean circuits to proof circuits

Input: $L_{\text{DC}}(\mathbf{C})$ for $\mathbf{C}$ a uniform family of polynomial-size and $\log^i$ -depth Boolean circuits over $\mathfrak{B}_u$.

Output: $L_{\text{DC}}(\mathbf{P})$ for $\mathbf{P}$ a uniform family of polynomial-size and $\log^i$ -depth proof circuits over $\mathcal{P}_u$, such that for all $n \in \mathbb{N}$, for all $\overrightarrow{b} \equiv b_{i_1}, \dots, b_{i_n}$, $P_n(\overrightarrow{b}) \rightarrow_{\text{ev.}} b_j$ iff $C_n(i_1, \dots, i_n)$ evaluates to j .

Theorem 2.3.3. *Problem C belongs to AC^0 .*

Proof. The translation from $\mathbf{C}$ to $\mathbf{P}$ is obvious, it relies on coding: for every n , a first constant-depth circuit associates to every gate of C_n the corresponding piece simulating its Boolean function, i.e. $\wedge^k$ is translated to $\mathcal{C}_{\text{onj}}^k$, $\vee^k$ by $\mathcal{D}_{\text{isj}}^k$ and $\neg$ by $\mathcal{N}_{\text{eg}}$. If the fan-out of this gate is $k > 1$, a $\mathcal{D}_{\text{upl}}^k$ piece is associated to the exit of the piece, and the pieces are connected as the gates. The input nodes are associated to the inputs of P_n . A second constant-depth circuit recollects the single free exit and the garbage of the pieces and connects them to the result tensor. The composition of these two Boolean circuits produces a constant-depth Boolean circuit that builds proof circuits in parallel.

It is easy to check that $\mathcal{C}_{\text{onj}}^k$, $\mathcal{D}_{\text{isj}}^k$ and $\mathcal{N}_{\text{eg}}^k$ represent $\wedge^k$, $\vee^k$ and $\neg$ respectively, that $\mathcal{D}_{\text{upl}}^k$ duplicates a value k times, and b_0 and b_1 represent 0 and 1 by convention. The composition of these pieces does not raise any trouble: P_n effectively simulates C_n on every input of size n .

Concerning the bounds: the longest path between an entry or a constant and the result tensor goes through at most $2 \times d(C_n)$ pieces and we know by [Corollary 2.3.2](#) that the increase of the depth is linear in the number of pieces crossed. We conclude that $d(P_n) \leq 2 \times 3 \times d(C_n)$ and we may also notice that P_n normalises in $\mathcal{O}(d(C_n))$ parallel steps.

Concerning the size, by simply counting we know that a gate of fan-in n and fan-out m is simulated by two pieces¹³ made of $\mathcal{O}(m+n)$ links. As the number of edges in C_n is bounded by $|C_n|^2$, the size of P_n is at most $\mathcal{O}(|C_n|^2)$. **Q.E.D.**

Corollary 2.3.4. *For all $i \in \mathbb{N}$*

$$\text{AC}^i \subseteq \text{PCC}^i$$

Note also that if the family $\mathbf{C}$ is over $\mathfrak{B}_b$, $\mathbf{P}$ is made of pieces of $\mathcal{P}_b$, and so $\text{NC}^i \subseteq \text{bPCC}^i$ for $i > 0$, as the translation is in AC^0 .

A Boolean circuit with unbounded (resp. bounded) arity of size s is translated to a proof circuit of size quadratic (resp. linear) in s , whereas TERUI considers only unbounded Boolean circuits and translate them with Boolean proof nets of size $\mathcal{O}(s^5)$. The previous translation [[107](#), Theorem 2, p. 8] was in $\mathbf{L}$ and so the inclusion could not be applied to sublogarithmic classes. Our translation—thanks mostly to the easier garbage collection—needs less computational power, is more clear and beside lowers the size of the Boolean proof nets obtained.

Simulation of Proof Circuits by Boolean Circuits

Now we will prove the “reverse” inclusion, i.e. that $\text{PCC}^i \subseteq \text{AC}^{i+1}$. This part will not be proved with a translation: we do not build a Boolean circuit family that computes as a proof circuit family. Rather, we define a Boolean circuit family that can *evaluate* any proof circuit, i.e. decide given an input and the Direct Connection Language of its family if it evaluates to b_0 or b_1 .

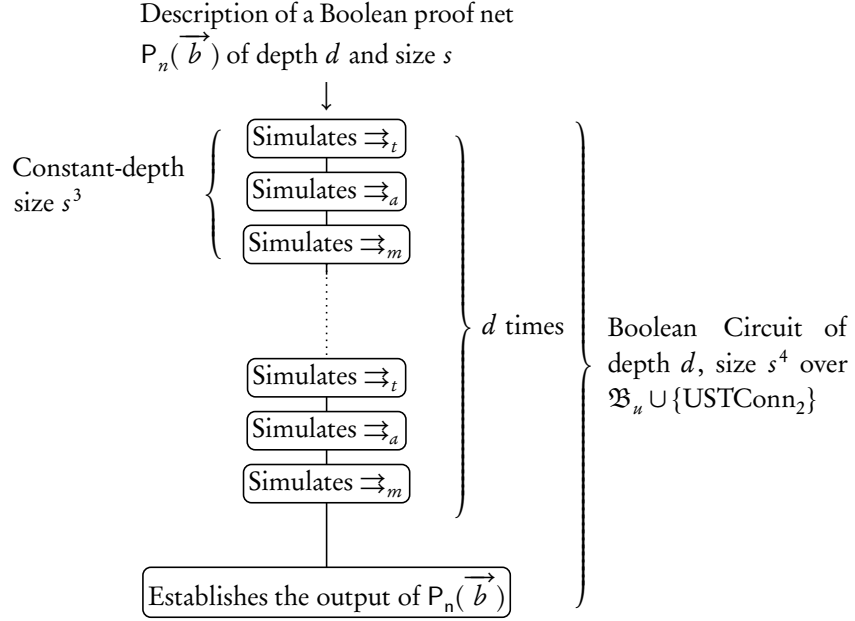
In other words, we define a family of Boolean circuits that is able to normalise in parallel any Boolean proof net, and so they need to perform t -reductions. This operation needs to identify maximal chains of axioms with cuts between them, and that require to solve the problem USTConn_2 ([Problem A](#)). We simply add USTConn_2 -gates to the basis of the Boolean circuits that will in constant depth identify the t -cuts and reduce them.

Problem D: Simulation from Boolean proof nets to Boolean circuits

Input: $L_{\text{DC}}(\mathbf{P})$ for $\mathbf{P}$ a polynomial-size, $\log^i$ -depth uniform Boolean proof net family.
Output: $L_{\text{DC}}(\mathbf{C})$ for $\mathbf{C}$ a uniform family of Boolean circuits $\mathbf{C}$ over $\mathfrak{B}_u \cup \{\text{USTConn}_2\}$,
 such that for all $\vec{b} \equiv b_{i_1} \dots b_{i_n}$, $P_n(\vec{b}) \rightarrow_{\text{ev.}} b_j$ iff $C_n(i_1, \dots, i_n)$ evaluates to j .

TERUI managed to perform this simulation thanks to *configurations*, another kind of description of the Boolean proof net that is updated during the course of the normalisation. We refer to [[130](#), [107](#)] for a technical presentation of this tool, and only give the sketch of the proof below. The [Figure 2.12](#) could help the reader to grasp the general structure of the Boolean circuit we are building.

¹³One for the computation, one for the duplication of the result.



Note that given a Boolean proof net in input, d steps of normalisation are enough to obtain the description of its normal form. The t -reduction needs USTConn_2 nodes to be made efficiently.

Figure 2.12: The simulation of a Boolean proof net by a Boolean circuit

For $r \in \{t, a, m\}$, an unbounded fan-in constant-depth Boolean circuit with $\mathcal{O}(|P_n|^3)$ gates—with USTConn_2 gates to identify chains of axioms if $r = t$ —is able to reduce all the r -cuts of P_n in parallel.

A first constant-depth circuit establishes the configuration from $L_{\text{DC}}(P)$ and constant-depth circuits update this configuration after steps of normalisation. Once the configuration of the normal form of P_n is obtained, a last constant-depth circuit identifies the first proof net connected to the result tensor and establishes if it is b_0 or b_1 —that is to say if the result of the evaluation is *false* or *true*.

As all the circuits are of constant depth, the depth of C_n is linear in $d(P_n)$. The size of C_n is $\mathcal{O}(|P_n|^4)$: every circuit simulating a parallel reduction needs $\mathcal{O}(|P_n|^3)$ gates and in the worst case—if $d(P_n)$ is linear in the size of the proof circuit— $\mathcal{O}(|P_n|)$ steps are needed to obtain the configuration of the cut-free Boolean proof net.

Theorem 2.3.5 ([107, Theorem 5, p. 9]). *Problem D is in L.*

Corollary 2.3.6. $\text{PCC}^i \subseteq \text{mBN}^i \subseteq \text{AC}^{i+1}$ for all i .

Proof. We know since [136] that USTConn is L-complete, but there is no proof that a weaker problem as USTConn_2 belongs to a weaker class of complexity. But since $\text{L} \subseteq \text{AC}^1$, we could replace all USTConn_2 gates by log-depth Boolean circuits, thus yielding $\text{AC}^i(\text{USTConn}_2) \subseteq \text{AC}^{i+1}$. The previous theorem and Theorem 2.2.3 complete the chain of inclusion. Q.E.D.

The simulation is slightly different from the translation: the Boolean circuit does not have to identify the pieces or any mechanism of computation of P_n , but simply to apply $\Rightarrow_t$, $\Rightarrow_a$ and $\Rightarrow_m$ until it reaches a normal form and then look at the obtained value.

Remark. We have focused on sublogarithmic classes of complexity, but we can draw more general conclusions by re-introducing USTConn_2 . TERUI proved that USTConn_2 could be represented by Boolean proof nets of constant-depth and polynomial size. So we could define “ USTConn_2 -piece”, $\mathcal{P}_{u+s} = \mathcal{P}_u \cup \{\text{USTConn}_2\}$ and $s + \text{PCC}^i$ as the class of languages recognized by $\log^i$ -depth, poly-size proof circuits over $\mathcal{P}_{u+s}$. It would be easy to conclude that for all $i \in \mathbb{N}$, $s + \text{PCC}^i = \text{AC}^i(\text{USTConn}_2) = \text{mBN}^i$.

2.4 Correspondence With Alternating Turing Machines

Alternating Turing Machines (ATMs) are “the” historical way of characterizing parallel computation. COOK [27], in 1985, defines AC^k as “the class of all problems solvable by an ATM in space $\mathcal{O}(\log n)$ and alternation depth $\mathcal{O}(\log^k n)$.” Remark by the way that the “A” of AC stands for “alternation”.

In this setting, there is no worries to have regarding uniformity, but some slight changes will be made in order to define “families of computation graphs” for ATMs (Definition 2.4.2). That will require that all the runs of an ATM on inputs of the same size have the same “shape” (Definition 2.4.1). This will help us to get closer to the computational behaviour of the proof circuits and to establish bridges.

This other characterization of the computational power of proof circuits has several advantages:

- it helps to grasp in a different manner parallel computation,
- it makes uses of log-space computation, and helps us to introduce some oddities of this calculus,
- it uses the paradigm “ATMs as graphs” that we sketched in Section 1.1, for instance in Definition 1.1.5
- and, last but not least, it refines the bounds we proposed in the previous section.

Recall that we proved, for all $i \in \mathbb{N}$,

$$\text{AC}^i \subseteq \text{PCC}^i \subseteq \text{AC}^{i+1} \qquad \text{NC}^i \subseteq \text{bPCC}^i \subseteq \text{PCC}^i \subseteq \text{bPCC}^{i+1}$$

Here we will prove¹⁴ for $j > 1$ and $k > 0$,

$$\text{NC}^j \subseteq \text{PCC}^j \qquad \text{bPCC}^k \subseteq \text{AC}^k$$

Whereas the first inclusion could be obtained from the previous section, the second inclusion is a novelty and enlightens the position of bounded proof circuits.

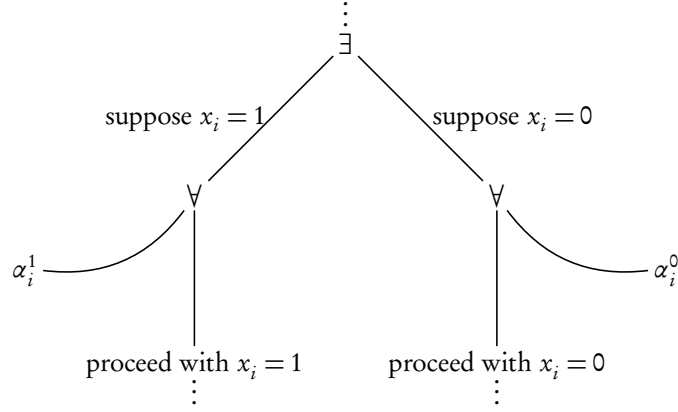
Unfortunately, the simulation and translation we will define in the following are quite expansive, and those results do not apply to constant-depth circuits. Anyway, by combining the results of the previous section and the results¹⁵ to come, we obtain for all $i > 1$,

$$\text{NC}^i \subseteq \text{bPCC}^i \subseteq \text{AC}^i \subseteq \text{PCC}^i \subseteq \text{bPCC}^{i+1}$$

The results in this section are unpublished and new, they are the result of a strong motivation to link together three classical models of parallel calculus. We begin with some remarks on ATMs, we

¹⁴See errata.

¹⁵See errata.



Rather than directly reading x_i , the ATM pursue the computation in parallel: on one way it makes as if x_i was 0, on the other it makes as if it was 1. The state α_i^j leads to acceptance iff $x_i = j$, so it discriminate between the correct and the wrong guess.

Figure 2.13: The access to the input in an input normal form ATM

will refine them to obtain more easily the results we expect. We take the liberty to represent the ATM program as a graph, following Definition 1.1.5.

Definition 2.4.1 (Input normal form). An ATM is in *input normal form*,¹⁶ if for every input $x \equiv x_1 \dots x_n$, for all $1 \leq i \leq n$, all accesses of M to x_i are of the form illustrated in Figure 2.13.

This mechanism is the core of the parallelism of ATMs: here, thanks to the random access to the input, checking if $x_i = j$ has a constant cost, but imagine that we have to perform a computation with a significant cost in terms of space, time or alternation. With this “guess and check” routine, we can safely perform this checking and the course of the computation *in parallel*. Moreover, it is possible to guess values longer than a single bit (cf. Definition 2.4.4).

We will massively use this procedure in the following, but we dwell for now on one of its other advantages: the course of the computation is now relative to the size of the input. This comes from the remark that if M an ATM is in input normal forms and $|x| = |y|$, then $G_{M(x)} = G_{M(y)}$. The difference will rest on the labelling of the graph (recall Definition 1.1.3), leading to acceptance or rejection.

Definition 2.4.2 (Family of computation graphs). Given M an ATM, its *family of computation graphs* $G_M = (G_{M,n})_{n \in \mathbb{N}}$ is such that for every $p \in \mathbb{N}$, $G_{M,p}$ is the computation graph of M on input of size p . We will write $\alpha \vdash_{M,n} \beta$ if there is an edge between α and β in $G_{M,n}$.

The computation graph associated to M on input of size $|x|$, $G_{M,|x|}$, is of depth at most $\mathcal{O}(t(|x|))$ if M is bounded in time by t .¹⁷ If M is bounded in space by s , the number of configurations of M is bounded by $\mathcal{O}(2^{s(|x|)})$ and so is the number of nodes of $G_{M,|x|}$. The bounds on alternation has no consequence on the shape of the graph.

¹⁶This suppose that our ATMs have random access to their inputs [135, p. 58].

¹⁷This comes from the simple remark that it is always possible to halt the computation after $\mathcal{O}(t(|x|))$ transitions.

The *depth* of a node is the length of the computation path from the initial configuration to the configuration corresponding to this node. It will be the equivalent of the height of a piece as defined in Definition 2.2.6.

Translation from Alternating Turing Machines to Proof Circuits

Definition 2.4.3 (Translation, inspired from VOLLMER [135, p. 63]). We define inductively a translation $(\cdot)$ from a family of computation graphs G_M to a family of proof circuit $\mathbf{P}$.

Given a size of input n , we proceed as follow:

- For $1 \leq j \leq n$, there exists $a \in \mathbb{N}$ such that the input x_j is tested a times. If $a = 0$, we simply connect an edge labelled with x_j to a garbage port of the result tensor, otherwise we display a $\mathcal{D}_{\text{upl}}^a$ piece whose entry is labelled with x_j .
- We then consider the sinks (nodes of height 0), which are final configurations or tests. Suppose that the configuration is α and there exists $\beta_1, \dots, \beta_i$ such that for $1 \leq k \leq i$, $\beta_k \vdash_{M,n} \alpha$.
 - If $g(\alpha) = 0$ (resp. 1), the node is translated with b_0 (resp. b_1),
 - if α tests the if the j th bit of the entry is 0 (resp. 1),¹⁸ the node is translated with the piece $\mathcal{N}_{\text{eg}}$ connected (resp. by directly connecting it) to one of the free output of the $\mathcal{D}_{\text{upl}}^a$ piece connected to x_j .

In both cases the output of the piece is connected to $\mathcal{D}_{\text{upl}}^i$.

- If the height of the node is $h \geq 1$, we proceed as follow: suppose the configuration associated is α . There exists $j, k \in \mathbb{N}^*$, $\beta_1, \dots, \beta_j$ such that for all $1 \leq m \leq j$, $\beta_m \vdash_{M,n} \alpha$, and $\gamma_1, \dots, \gamma_k$ such that for $1 \leq l \leq k$, $\alpha \vdash_{M,n} \gamma_l$. By induction, we already know pieces translating $\gamma_1, \dots, \gamma_k$, because their height is strictly inferior to h . We know that $g(\alpha) \notin \{0, 1\}$ since α is not final.

If $g(\alpha) = \exists$ (resp. $\forall$), the node associated to α is $\mathcal{D}_{\text{isj}}^k$ (resp. $\mathcal{C}_{\text{onj}}^k$), with its k input connected to the translations of $\gamma_1, \dots, \gamma_k$. Its output is connected to $\mathcal{D}_{\text{upl}}^j$, which transmits the value to the pieces translating $\beta_1, \dots, \beta_j$.

We then add a result tensor connected to the output of the translation of the initial configuration and to the garbage. Since every computation graph is finite and presents no cycle, this translation ends. As we simply connected together pieces and added a result tensor, the structure obtained is a proof circuit—hence a proof net—and it has n edges labelled with $x_1, \dots, x_n$.

The Figure 2.14 might help the reader to check how simple this translation is. However, even if this algorithm is simple to understand, it has a log-space complexity, as shown below.

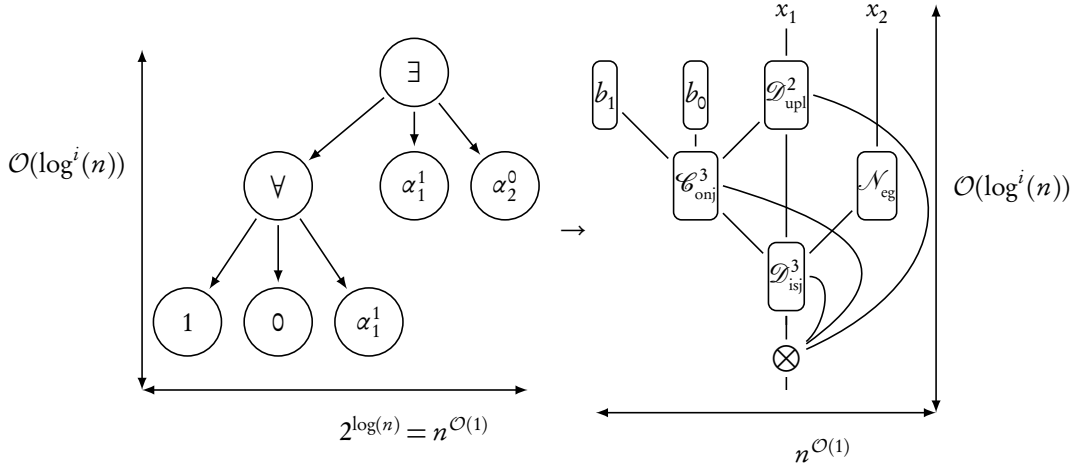
Problem E: Translation from ATM to proof circuits

Input: A log-space, $\log^i$ -time bounded ATM M .

Output: The Direct Connection Language of a uniform family of $\log^i$ -depth proof circuits $\mathbf{C}$ such that for all $n \in \mathbb{N}$, for all $\vec{b} \equiv b_{i_1}, \dots, b_{i_n}$, $C_n(\vec{b}) \rightarrow_{\text{ev.}} b_1$ iff $M(i_1, \dots, i_n)$ accepts.

Theorem 2.4.1. *Problem E belongs to L.*

¹⁸Something we wrote α_i^0 (resp. α_i^1) in Figure 2.13.



We take the Alternating Turing Machine here to be log-space and $\log^i$ -time bounded, and only specify the state type of every configuration (recall [Definition 1.1.2](#)). The configuration α_i^j is labelled with 1 iff $x_i = j$, 0 elsewhere. For convenience we represented the edges of the proof circuit from top to bottom, but of course the pieces are connected with cuts.

Figure 2.14: The translation from Alternating Turing Machines to proof circuits

Proof. First, remark that we can always suppose that the ATM is in input normal form: if it is not, there exists an ATM that accepts the same language with the same bounds.¹⁹ Secondly, remember that we sketched in [Section 1.1](#) how a deterministic Turing Machine with logarithmic space could output the computation graph $G_{M(i_1, \dots, i_n)}$ of a Turing Machine M . From that, it is easy to output with logarithmic resources $\widehat{G_{M,n}}$ a proof circuit that will reduce as expected.

Concerning the bounds: we know that the number of configurations of $M(x)$ —and so the size of $G_{M,n}$ — is bounded by $2^{O(\log(n))}$. Every configuration with m predecessors and k successors is translated in $\widehat{G_{M,n}}$ with a piece of size $O(m+k)$ and a number linear in n of $\mathcal{D}_{\text{upl}}$ piece may be added to duplicate the inputs. As $|\widehat{G_{M,n}}| = O(|G_{M,n}|) = 2^{O(\log(n))} = n^{O(1)}$, $|\widehat{G_{M,n}}|$ is polynomial in n .

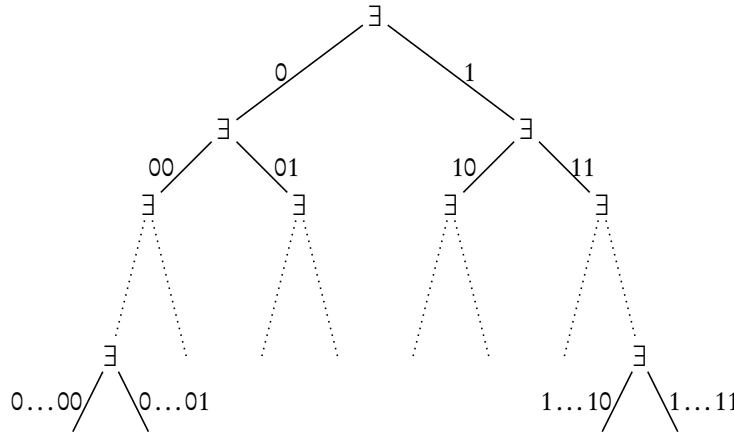
Regarding the depth, the longest computational path from the initial configuration to a final configuration is of length at most $\log^i(n)$. Every node is translated with at worst two pieces, so the depth of $\widehat{G_{M,n}}$ is linear in $\log^i(n)$ by [Lemma 2.3.1](#) Q.E.D.

Corollary 2.4.2. *For all $i > 1$,*

$$\text{STA}(\log, \log^i, *) \subseteq \text{PCC}^i$$

Unfortunately this result could not be extended to the sublogarithmic case, as the reduction developed in the previous proof belongs to **L**. This result concerns only $\text{NC}^2 = \text{STA}(\log, \log^2, *)$ and above, which has more resources than **L**.

¹⁹Recall that being in input normal form has a cost in alternation only, but we say nothing regarding this bound in this theorem.



The tree has a depth equal to the size of the guessed value.

Figure 2.15: How a value is guessed by an Alternating Turing Machine

Simulation of Proof Circuits by an Alternating Turing Machine

We will perform here a typical log-space-computation: the size of a Boolean proof net is polynomial in the size of its input, so the Alternating Turing Machine cannot rewrite the description of the Boolean proof net until it reaches a cut-free form, because it simply does not have enough space. Instead, the Alternating Turing Machine will parse, or peruse, its Direct Connection Language in order to find directly the result of the normalisation. We save space by doing that way and are able to write down an algorithm *à la* log-space. This procedure shares a lot with the pointer machinery, a topic that will be the center of Chapter 4.

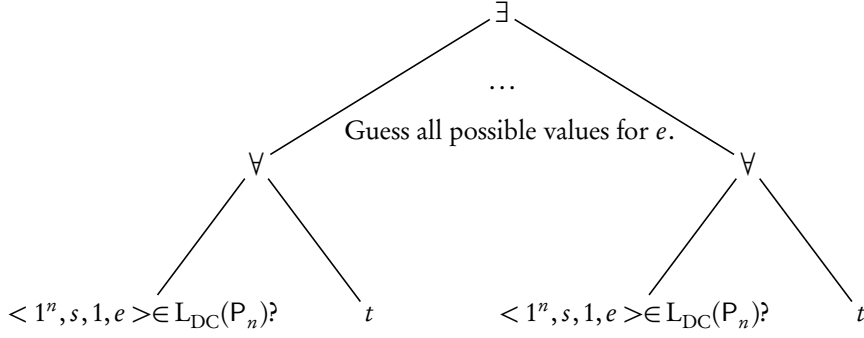
It will rely at the same time on the Direct Connection Language of the Boolean proof nets (recall Definition 2.2.3), on their uniformity, and for the last part on the constrained structure of a proof circuit. We will use the same mechanism as in the guessing of the input by an input normal form ATM to guess information regarding the proof circuit we are trying to normalise.

Definition 2.4.4 (How to guess?). We show how to guess numbers of links; to guess other values of constant size is essentially the same mechanism. Given an input $\vec{b}$ of size n , it takes place $\log(|P_n|)$ to write the number of a link in binary, and as $|P_n| = n^{\mathcal{O}(1)}$, that is to say $\mathcal{O}(\log(n))$.

The semantic of the guessing is the following : “there is a 1st bit which can be 0 or 1, there is a 2nd bit which can be 0 or 1, ..., there is a i th bit which can be 0 or 1”. The Figure 2.15 shows how it works. It takes time i , space i and has a constant cost in alternation to guess a value of size i .

Once the value we are looking for is guessed, we have to check that it is the *actual* value that we are looking for. In other terms, if we are looking for the number of a link e such that e is connected to the first auxiliary port of the link number s in the Boolean proof net P_n , we have to make sure that $\langle 1^n, s, 1, e \rangle \in L_{DC}(P_n)$. This question —once e is guessed— takes time $\mathcal{O}(\log(n))$ on a deterministic Turing Machine to be answered, since our proof circuits are uniform.

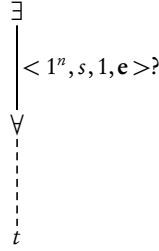
We present in Figure 2.16 the composition of the guessing and of the checking by the uniformity machine. Remark that an important amount of useless work is done: this is typical of a log-space



We guess all possible values for e by guessing all the binary string of length $\log(|P_n|)$. Then, for every of them, we test if it is the value we are looking for, by asking the uniformity machine if $\langle 1^n, s, 1, e \rangle \in L_{DC}(P_n)$.

Among all the values guessed for e , only one will satisfy this belonging, for the first auxiliary port of the link numbered e in P_n is connected to only one other link. So all the branches except one will reject, and the computation will go on with the actual value for e .

We will write in Figure 2.17 this routine as



Where e is in bold because it is the value we are looking for, and a dashed line means that the uniformity machine is been asked. A constant number of values may be guessed and checked at the same time, with only one alternation, as long as there is enough space to store the values.

Figure 2.16: A shorthand for the “guess and check” routine

computation. As we do not have enough space to memorize the number of the links, we just calculate them again and again, when we need them.

Now that we defined how to obtain with efficiency some information regarding a Boolean proof net, we have to show how an ATM can use them wisely to establish the outcome of its normalisation. To try to establish an algorithm with logarithmic resources for any Boolean proof net is a bit naive, for the complexity of cut-elimination is quite high: it has been proven by MAIRSON and TERUI [100] that the cut elimination procedure is $\mathbf{P}$ -complete for IMLL—the intuitionist fragment of MLL. As they write it [100, p. 24], “This result contradicts folkloric intuitions that **MLL** proofnets could be normalized in logarithmic space—that is, with only a finite number of pointers into the proofnet, presumably following paths in the style of the geometry of interaction.”

We proved it anyhow for a constrained fragment of the proof nets of $\mathbf{MLL}_u$: the proof circuits. In this setting, we always know “what to expects”, as proof circuits are built following the simple pattern of composing pieces. To each entry and exit of a piece, we may associate a value, i.e. the edges labelled e and s in Table 2.8 may be typed with $\mathbf{B}$, and they will “transport” a boolean value. Our algorithm will find for every piece, starting from the exit of the proof circuit, what are the values carried by those

edges. We cannot associate to “the value carried is 0” rejection and to “the value carried is 1” acceptance, because we rely on tests, on guess and checks, where rejection means “you did not made the right hypothesis, I am not even able to perform this computation”.

So we define a recursive function f that given the number of a link and an expected value, leads to acceptance if and only if the edge connected to the principal port of that link will “transport” the expected value. That distinguishes “rejection cause the test fails” and “rejection because the value is 0”.

Definition 2.4.5 ($f : \text{number of a link} \times \text{value expected} \rightarrow \{\text{accept}, \text{reject}\}$). Given s the number of a link and $v \in \{0, 1\}$, f will accept iff the proof net connected to the principal port of s will normalise to b_v . This function is recursive and we give in Figure 2.17 a representation of the computational graph of an ATM computing this function: it massively uses the “guess and check” procedure and the shorthand we introduced in Figure 2.16. It does not cover the case when s is connected to an input: in this case, simply accept if this input is b_v , reject elsewhere.

We give some explanation on how this function computes: it takes as input $v \in \{0, 1\}$ —the “expected” value— and s —the number of a link, that may safely be identified with the edge connected to its principal port, i.e. with the edge labelled s in Table 2.8, the exit of a piece.

Then f proceeds by cases, with $\bar{v} = 0$ if $v = 1$, 0 elsewhere:

1. If s is the output of a piece b_i ($i \in \{0, 1\}$) accepts if $i = v$, rejects elsewhere.
2. If s is the output of a piece $\mathcal{N}_{\text{eg}}$ whose entry is connected to the link number e , do as $f(\bar{v}, e)$.
3. If s is the output of a piece $\mathcal{D}_{\text{upl}}^i$ whose entry is connected to the link number e , do as $f(v, e)$.²⁰
4. If s is the output of a piece $\mathcal{C}_{\text{onj}}^2$ whose entries are connected to the links number e_1 and e_2 , and if $v = 1$ (resp. $v = 0$), universally (resp. existentially) branch on $f(v, e_1)$ and $f(v, e_2)$.
5. If s is the output of a piece $\mathcal{D}_{\text{isj}}^2$ whose entries are connected to the links number e_1 and e_2 , and if if $v = 1$ (resp. $v = 0$), existentially (resp. universally) branch on $f(v, e_1)$ and $f(v, e_2)$.

We prove that this procedure ends thanks to an inductive proof on the depth and by using Lemma 2.3.1: in the $\mathcal{N}_{\text{eg}}$ case, the call to f is performed on a link at same depth. In the $\mathcal{D}_{\text{upl}}^i$ case, the call is performed at a greater depth. In the $\mathcal{C}_{\text{onj}}^2$ and $\mathcal{D}_{\text{upl}}^2$ cases, two calls are made: one at the same depth, the other at superior depth. In the b_i case, no recursive call is made.

As the $\mathcal{N}_{\text{eg}}$ pieces are limited to the entries (recall Definition 2.2.5), we can consider this case only as a constant factor. Except for this case, each iteration performs at least one call to f with a link of depth superior to the previous call. So the procedure ends, and in the worst case $\mathcal{O}(d(P_{|n|}))$ calls are performed. That f leads to accept with input v , s iff the link v will transport the value v is just a matter of common sense.

The following theorem and its proof are wrong, see errata.

Theorem 2.4.3. For all $i > 0$.

$$\text{bPCC}^i \subseteq \text{STA}(\log, *, \log^i)$$

Proof. We build our ATM M as follows:

²⁰A remark should be made about this treatment of the duplication: there is no sharing. Every time we want to know the output of a duplication, we test its input. So to say, the circuit is “splited” and treated as it was a tree.

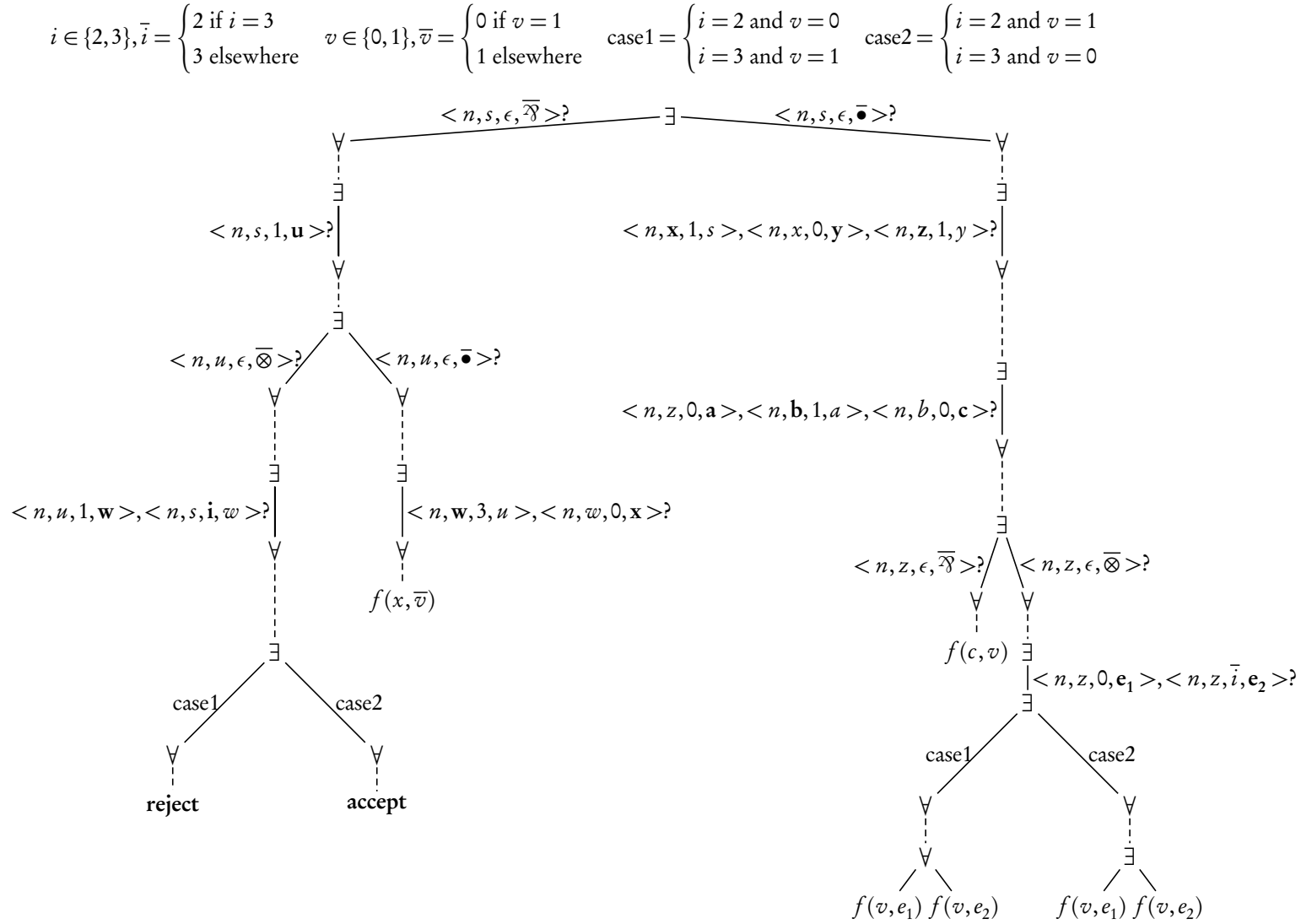


Figure 2.17: A function to compute the value carried by an edge

- It takes as input $L_{DC}(\mathbf{P})$ for $\mathbf{P}$ a uniform, polynomial-size, $\log^i$ -depth proof circuit family over $\mathcal{P}_b$, and an input $\vec{b} = i_1, \dots, i_n$.
- It identifies the number of the result tensor of P_n .²¹
- It guesses the number of the link connected to the first auxiliary port of the result tensor, and it runs f with this value and 1.

It should be clear that M accepts iff $P_n(\vec{b}) \rightarrow_{ev.} b_1$, and stops rejecting elsewhere.

Regarding the bounds:

Space At every moment M has only to store a constant number of addresses of links. We can at each transition forget every data except the expected value and the number of the link under treatment. A logarithmic amount of space is sufficient to store those values.

Alternation Each call to f uses a constant number of alternations, and $\mathcal{O}(d(P_n))$ calls are made. As P_n is of depth $\log^i(n)$, $\mathcal{O}(\log^i(n))$ calls are made.

Q.E.D.

Two remarks should be done:

- This algorithm cannot be extended as it to Boolean proof nets, even if they are uniform. As our pieces are “normed”, the ATM knows what to expect, and so we can encode in the algorithm itself the behaviour to adopt. A Boolean proof net in all generality is a too complex object to be treated the same way.
- We work in the bounded case for complexity reasons. Our algorithm could work with just a slight modification in the unbounded case, but the bounds would be problematic. We would need to introduce a routine that guess and checks the arity of every $\mathcal{D}_{isj}$ and $\mathcal{C}_{onj}$ pieces. But that would re-introduce linearity in our algorithm, for the ATM needs to establish the number of the $\otimes$ corresponding to the $i + 1$ th entry before establishing the number of the $\otimes$ corresponding to the i th entry. As in the worst case, the fan-in of the piece is relative to the size of the proof circuit, we obtain a bound on time relative to the *size* of our Boolean proof net.

Results and perspectives

What have we done?

We recalled a reliable way of representing Boolean values and Boolean functions with MLL_u proof nets. We then restated a depth-efficient way of computing with them, letting them able to be normalised by Boolean circuits, thus proving (Corollary 2.3.6), for all $i \in \mathbb{N}$,

$$PCC^i \subseteq AC^{i+1}$$

²¹That is to say to find the only $\otimes$ -link with its principal port unconnected. That can be done in parallel, or more simply we can by convention says that the first link described in the Direct Connection Language of the proof net is the result tensor.

We then constrained the structure of the Boolean proof nets by introducing proof circuits and gave an AC^0 -reduction from them to Boolean circuits, establishing —even when $i = 0$ (Corollary 2.3.4)

$$AC^i \subseteq PCC^i$$

We adapted the classical representation of Alternating Turing Machines in Boolean circuits toward proof circuits, and that gave us (Corollary 2.4.2), a new proof that for $k > 1$

$$NC^k \subseteq PCC^k$$

At last, we deeply used the parallelism of Alternating Turing Machines to prove that they were able with log-space resources to obtain the cut-free form of any proof circuits with bounded fan-in. By combining Theorem 2.4.3 with the equality $AC^i = STA(\log, *, \log^i)$ for all $i \geq 1$, we obtain²², for $k > 1$

$$bPCC^k \subseteq AC^k$$

We managed to lighten the simulation of the Boolean functions by proof nets and extended previous results to sublogarithmic classes of complexity. We also managed to build strong gateways between Boolean proof nets and ATMs, and even if those results does not apply to the very small classes, it shed a new light on the links between logarithmic space computation and proof nets.

What could be done?

This work could be pushed further in several directions:

1. GIRARD [49, p. 8] was very enthusiast regarding non-commutative Linear Logic, “which should play a prominent role in the future”. This fragment does not forbid exchange, but constrains them by allowing only circular permutations, that is to say planar proof nets. This particular framework led to numerous fruitful studies, including specific correctness criterion [103]. Her, our modelization of computation used permutations, something that was at a hidden level in the sequent calculus. We made explicit those exchanges thanks to the proof nets, but other approaches has been developed, for instance by ANDREOLI, MAIELI, and RUET [3]. Although the main interest of our setting is to compute in parallel, it could be interesting to see how computation may be represented by cyclic linear logic, which has a part of non-determinism in it.
2. Non-determinism was also simulated with the Boolean proof nets by introducing a part of the additive fragment by MOGBIL [105]. The simplicity of the proof circuits allows to write quite directly “by hands” proof nets and to compute with them. By introducing small doses of non-determinism in the pieces, it could be possible to understand precisely what is at stake in this setting. Some promising results between proof nets and non-deterministic were obtained by MAUREL [101]. But non-deterministic computation in this setting looks more like a juxtaposition of deterministic computation than true non-determinism.
3. Semi-Unbounded-Fanin AC (SAC^k) is the class of decision problems solvable by a family of $\log^k(n)$ -depth circuits with unbounded-fanin $\vee$ and bounded-fanin $\wedge$ gates. Of course, we take as a constraint that negations are only allowed at the input level. It was introduced by BORODIN et al. [20] and

²²See errata.

proven that this class was closed under complement for every $k > 0$. We know also that $\mathbf{SAC}^1 = \mathbf{LOGCFL/poly}$ since the work of VENKATESWARAN [134]. To investigate in this direction could offer a very nice framework to grasp non-uniform classes of complexity, recently sketched by MAZZA [102].

Computing in the Hyperfinite Factor

Contents

2.1	Basic Definitions: Boolean Circuits and MLL_u	22
	Boolean circuits	
	Unbounded Multiplicative Linear Logic	
	Proof Nets for MLL_u	
2.2	Boolean Proof Nets	28
	Proof Nets that compute Boolean functions	
	Proof Circuits	
	Contractibility criterion: Proof Circuits are Boolean Proof Nets	
2.3	Correspondence With Boolean Circuits	41
	Translation from Boolean Circuits to Proof Circuits	
	Simulation of Proof Circuits by Boolean Circuits	
2.4	Correspondence With Alternating Turing Machines	45
	Translation from Alternating Turing Machines to Proof Circuits	
	Simulation of Proof Circuits by an Alternating Turing Machine	

IN this chapter, we are going to define another model of computation inspired from the Linear Logic, and more precisely from the Geometry of Interaction (GoI). Although the inspiration comes from Linear Logic, the objects we are going to manipulate are purely mathematical. This comes from the GoI program, and more precisely from its fifth version[57], that tells us that operators in the Π_1 hyperfinite von Neumann algebra describe totally the computational behaviour of proofs. Others variants of the GoI programs were already used with complexity concerns, and we present some of them in Section 3.1, before making clear that our approach here is a total novelty.

We build patiently in Section 3.3 the representation of binary words with matrices starting from proofs, and the ELL representation in particular. From this construction, it is rather easy to build a representation of integers in the hyperfinite factor, and then to define programs that will interact with

them (or, as we call them later, *observations*), in [Section 3.4](#). We take the time in the second subsection of [Section 3.4](#) to explain and motivate GIRARD’s choices, and to dwell on the difference of the approach we developed with Thomas SEILLER.

Once the representation of computation in this setting is defined, we fix a parameter (the *normative pair*) and constraint the program on a mathematical level, with the positiveness of the coefficients and the 1-norm. We then define a particular normative pair, which relies heavily on permutations. The group of finite permutations over $\mathbb{N}$ is well-known to the Linear Logic community for it was used to interpret the multiplicative fragment of Linear Logic.

The last part is devoted to prove that the computation that takes place in the Hyperfinite factor can be grasped with logarithmic space. This can be achieved thanks to a complex lemma presented in [Section 3.5](#), whose purpose can easily be explained as the trip back from hyperfinite factor to finite matrices. It is the first step toward a *mathematical* characterisation of the log-space computation. It borrows some concepts to the Implicit Computational Complexity (ICC) approach, but we will stick in this chapter to a mathematical level.

This setting was first proposed by GIRARD [59] and enhanced by AUBERT and SEILLER[10, 11] recently. This chapter borrows a large part of the content of those works, but presents it in a innovative ways, with examples and a more pedagogical development (see in particular [Section 3.2](#), when the construction is explained step-by-step for the unary representation of the integers).

This chapter is to be read at the light of its dual [Chapter 4](#), where the computational behaviour of the observations in a particular normative pair is compared to pointer machines. It is largely in debt to SEILLER’s precious intuitions, his patience and well understanding of the von Neumann algebras. Even if it uses some quite complex mathematical tools, as the von Neumann algebra or the crossed-product, we preferred to define those technical tools in the [Appendix A](#). The technical details are not what is important in this innovative construction, so we felt free to give rather global ideas to explain it. Yet, some useful element are reminded in the second moment of [Section 3.1](#).

3.1 Complexity and Geometry of Interaction

Past approaches

Linear Logic always had a tight link to complexity, thanks to its resource-awareness. Geometry of Interaction, while giving a dynamics semantics to Linear Logic, preserved that concern and gave new tools to express cut-elimination. The famous “execution formula” [48, p. 222] defined a way to express strong normalisation as the nilpotency of a set of operators.

This first approach gave rise to numerous semantics for various fragments of Linear Logic, much of them being sensitive to the cost of normalization in terms of time: the studies of IMELL [33] and ELL [14] being probably the two most famous examples. This first paper use some games and ICC techniques to build an interpretation for IMELL proofs that reflects their time-complexity. The second paper is a bit closer to our work, for it uses an algebra of clauses to interpret proofs of ELL. In this setting, proof nets are represented with operators, and a novel way of executing them—independent from the nilpotency—is defined and proven to preserve some complexity properties. Context semantics is one of the model of the GoI and was used to prove some interesting computational complexity properties [31].

On a more syntactical level (even if the GoI tries to “scramble” the distinction between syntax and semantics), the GoI was often understood as a novel approach that subsumes the proof nets into a set

of permutations. Stated differently, the type of the proof net is not the center of our object anymore: all the information is contained in the *paths*. This perspective gave rise to some hope that proof nets could be parsed with pointers and described with log-space (but we saw in the previous chapter, page 50, how MAIRSON and TERUI [100] proved that it was only “a rumour”, provided of course that $L \neq P$). Anyway, this gave numerous works where Linear Logic was linked to complexity and apprehended thanks to various kind of token machines.

Last but not least, the bridges between Linear Logic and operator algebra led to think that Linear Logic could be a place to express quantum computation. Indeed, quantum mechanics is expressed by mathematical formulations in a complex separable Hilbert space¹—as for the GoI. Since the pioneer work by SELINGER [122] that linked λ -calculus and quantum computation,² numerous attempts have been made to link those two world by the entanglement of von Neumann algebras, GIRARD [54] being the first. Quantum complexity was not left aside, and quantum circuits were proven to have a closed link with Linear Logic by DAL LAGO and FAGGIAN [32].

Complexity with operator algebra

We define in this chapter a new way of linking Linear Logic and complexity *via* von Neumann algebras by using operators as representations of algorithms. The multiplication of the representation of an integer and the representations of a program will be the execution of the program, so there is no need to define transition functions. Despite its high degree of mathematical abstraction, the model of computation we are going to define is a low-level one: every operation can be understood in its most modest details.

This attempt can be considered as the first success toward a purely mathematical way of defining a model of computation and complexity classes³ It uses operators acting on an infinite-dimensional (separable) Hilbert space, but can largely be understood in terms of matrices. It does not refer to any exterior measure of complexity and takes its place in the fast-growing field of ICC.

We give below some intuitions regarding the choice of those mathematical tools. The reader might refer to [Appendix A](#) for a short reminder of the algebraic and topological elements needed to grasp the usefull part of von Neumann algebras. It also contains some useful references and bibliographical indications. We simply give in the following some insights of the mathematical properties of the von Neumann algebras. We will not explore the definitions and properties of von Neumann algebras, factors, hyperfiniteness and crossed-product in this chapter. We use only the core of those notions in our construction, but they do not really play an important role. Although those tools won't be used before [Section 3.4](#), we remind of them now, so that we won't have to interrupt the quite subtle construction that will follows.

We begin by recalling that a Banach Algebra—a $\mathbb{C}$ -vector space enriched with the structure of an algebra, a norm, and complete for that norm—endowed with an involution is a C^* -algebra. A von Neumann algebra $\mathfrak{M}$ is a C^* -algebra of operators⁴ that can be understood—depending on the reader's preferences—as

¹This is not the place to develop such a complex subject, but the interested reader may find some useful hints in the textbook of NIELSEN and CHUANG [110].

²One may also have a look at another pionner work [133, Section 9].

³Ugo DAL LAGO kindly made me notice that “any model of computation is defined mathematically”, so that this was a bit pretentious. Yet, it seems to us that the framework we are going to define can be fully understood from a mathematical point of view, in an innovative way.

⁴An operator is a continuous or equivalently bounded linear maps between Hilbert spaces. The set of operators on a Hilbert space $\mathbb{H}$ will be denoted $\mathcal{B}(\mathbb{H})$.

a C^* -subalgebra of $\mathcal{B}(\mathbb{H})$, which is closed for one of the topology among the weak operator, the strong, the ultrastrong or the ultraweak topologies;

a C^* -algebra that have a predual, i.e. there exists a Banach algebra such that $\mathfrak{M}$ is its dual;

a C^* -algebra equal to its bi-commutant, i.e. the set of elements of $\mathcal{B}(\mathbb{H})$ which commute with the elements that commute with themselves.

The beauty of von Neumann algebras comes partially from this plurality of definitions, which allowed transfers from a mathematical field to another.

The *hyperfinite* von Neumann algebras are of particular interest to us, because we can approximate every of their operators with finite matrices. We can also quotient the von Neumann algebras to study only their *factors*, which can be classified according to their *projections*. For most types hyperfinite factors are unique, and it is the case for the factor we will use in the following: the type II_1 hyperfinite factor.

As we have already said, permutations are a natural way of interpreting the multiplicative fragment of Linear Logic. Yet, to express the power of exponentials, we need to be able to express infinite objects. Operator algebra is quite an obvious choice, for it generalizes finite permutation thanks to *partial isometries*. A partial isometry u establishes an isometric bijection between its domain—the closed subspace associated with the projection u^*u —and its image—the closed subspace associated with the projection uu^* .⁵ Partial isometries have a tight bound with *projections*, i.e. operators which give an orthonormal projection of $\mathbb{H}$ onto a closed subspace of $\mathbb{H}$. The permutation of the elements of an infinite separable Hilbert space generalise quite elegantly the permutations.

Permutations will be re-integrated in this setting on a second level, thanks to the crossed product. This operation internalizes the action of a group—in this setting, the group of finite permutations over $\mathbb{N}$ —into the von Neumann algebra. This operation won't modify our object—which will remain the II_1 hyperfinite factor—but it will allow us in the next chapter to study easily the computational behaviour of our object.

3.2 First Taste: Tallies as Matrices

We will begin by explaining how integers in the unary notation (“tally sticks”) can be represented as matrices *via* the Curry-Howard correspondence. It has two interests: to see the computation from a dynamics point of view, and to represent any proof (hence, integer) inside the same object. As the construction is quite complex, we felt it necessary to begin with the unary representation, simpler to grasp. We won't use it in the following, so the reader quite used to this idea may safely skip this section and go straightforward to [Section 3.3](#). Nevertheless, most of the explanations we make in the following are still valid when applied to the binary representation.

The representation of integers as matrices goes through the following journey:

1. Represent integers with λ -terms (Church encoding).
2. Represent λ -terms with derivations (Curry-Howard isomorphism).
3. Represent derivations with proof nets (“Parallel syntax”).

⁵A restatement of [58, p. 400].

$$\frac{\frac{\frac{\frac{\frac{\frac{\frac{\vdash X_2, X_3^\perp}{\text{ax.}}}{\vdash X_{m-1}, X_1^\perp}{\text{ax.}}}{\vdash X_1^\perp \otimes X_2, X_3^\perp, X_{m-1}}}{\vdash X_4, X_5^\perp}{\text{ax.}}}{\vdash X_1^\perp \otimes X_2, X_3^\perp \otimes X_4, X_5^\perp, X_{m-1}}}{\vdots}}{\frac{\frac{\frac{\vdash X_{m-2}, X_m^\perp}{\text{ax.}}}{\vdash X_1^\perp \otimes X_2, \dots, X_{m-3}^\perp, X_{m-1}}}{\vdash X_1^\perp \otimes X_2, \dots, X_{m-3}^\perp \otimes X_{m-2}, X_{m-1}, X_m^\perp}}{\vdash X_1^\perp \otimes X_2, \dots, X_{m-3}^\perp \otimes X_{m-2}, X_{m-1} \wp X_m^\perp} \wp \otimes$$

We give in all generality the derivation of the unary representation of $n \in \mathbb{N}^*$, remember that we set $2n + n = m$.

Figure 3.1: The derivation of a unary integer, without contraction rule.

4. Represent proof nets with matrices (Geometry of Interaction).

We will in the following make one step further: to embed those matrices in the type II_1 hyperfinite factor. Those unary integers are not efficient from a complexity point of view: they are introduced only for pedagogical reasons, to show how data can be represented by matrices. A first version of this encoding—more in the “GoI V-style”—was developed [57, section A.3, p. 31], and then simplified [59, p. 251]. The version we propose here rather comes from this latter simplification, but has the pretension to be a bit more explanatory. We will first present it without the contraction rule, and introduce this subtlety in the next subsection.

First representation of the unary integers, without the contraction rule

Without the contraction rule, $n \in \mathbb{N}$ can be represented as a proof of

$$\underbrace{(X \Rightarrow X), \dots, (X \Rightarrow X)}_{n \text{ times}} \vdash X \Rightarrow X$$

This representation is not exactly the Church encoding of numerals, that represent n as the higher-order function that maps any function f to $f \circ f \circ \dots \circ f$. Due to the absence of contraction, that could allow us to get sure that *the same* function is applied n times,⁶ this encoding should rather be seen as the function that takes n functions $f_1, \dots, f_n$ and returns the composition of them $f_n \circ \dots \circ f_1$, whose corresponding λ -term would be $\lambda f_1 \dots \lambda f_n \lambda x \cdot f_n(\dots(f_1(x))\dots)$. We know that they all have the same type $X \Rightarrow X$, but we have no insurance whatsoever that they are identical.

To represent $n \in \mathbb{N}$, this encoding uses $2n + 2$ times the literal X , and we set for convenience that $2n + 2 = m$. We add in the following a subscript to help us to differentiate the m different occurrences of X , but keep in mind that the one-sided Linear Logic proof that corresponds to n represented in Figure 3.1 contains a single literal.⁷ The bits are “chained” with $\otimes$, and $\wp$ represents the end of the chain: it will be at the same time its beginning and its end, thus yielding a circular representation. But we have time for that and should first remark that the poverty of this proof is even more clear if we represent it as a proof net, as in Figure 3.2.

⁶This approach will be developed in the following subsection. We felt it necessary to begin with the an even simpler case.

⁷Otherwise the axiom for instance introducing at the same time X_2 and $X_3^\perp$ would be false.

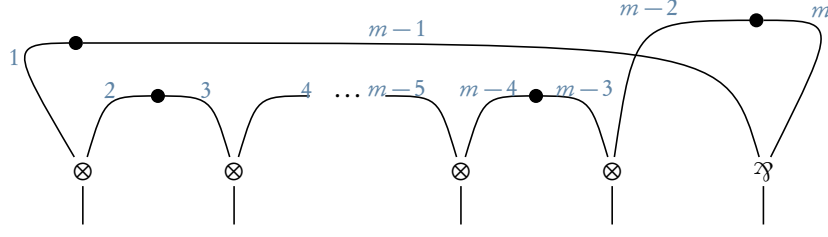


Figure 3.2: Unaries as proof net

Nevertheless, we can then represent for all $n \in \mathbb{N}^*$ its proof by a square matrix L_n ,⁸ with one column and one row per occurrence of X —i.e. of dimension $m \times m$ — in the corresponding proof. We interpret the axiom rule

$$\frac{}{\vdash X_j, X_i^\perp} \text{ax.}$$

as the symmetric exchange between i and j (i.e. $L_n[i, j] = 1$ iff there is an axiom between $X_i^\perp$ and X_j) and we define τ to be the anti-diagonal matrix of dimension 2. For all $n \in \mathbb{N}^*$, we interpret the proof corresponding to n with a $m \times m$ matrix:

$$\tau = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad L_n := \begin{pmatrix} 0 & 0 & 0 & \dots & 0 & 0 & 0 & 1 & 0 \\ 0 & & \tau & \dots & 0 & 0 & 0 & 0 & 0 \\ 0 & & & \dots & 0 & 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & & \vdots & \vdots \\ 0 & 0 & 0 & \dots & & \tau & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & & & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & \dots & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 1 & 0 & 0 \end{pmatrix}$$

Of course this representation is not satisfactory, for the dimension of the matrix itself —as for the type— tells us the represented integer. It inflates the space needed to encode the information without giving any dynamics. But this is the first step toward a more efficient representation. We add in the following the first enhancement: to allow the contractions.

Maybe we should before that remark that L_n is symmetric and that its diagonal is empty. Some more interesting properties are yet to come when we will have defined the associated projections, but we keep this further step for the following subsection.

Unary integers with the contraction rule

In λ -calculus, the Church representation yields that $n \in \mathbb{N}^*$ is encoded by

$$\lambda f \lambda x \cdot \underbrace{f(f(\dots(f x) \dots))}_{n \text{ times}}$$

⁸The notation we use regarding matrices is summed up in Definition A.2.5.

and the Curry-Howard correspondence tells us that the construction of this λ -term corresponds to the derivation of the type $(X \Rightarrow X) \Rightarrow (X \Rightarrow X)$.

To get this representation, we allow the contraction rule.⁹ Any proof of conclusion

$$\underbrace{(X \Rightarrow X), \dots, (X \Rightarrow X)}_{n \text{ times}} \vdash X \Rightarrow X$$

can easily be transformed into a proof of conclusion $(X \Rightarrow X) \vdash (X \Rightarrow X)$, where X occurs only four times. This representation is isomorphic to the representation of the integers in system F,¹⁰ which associate to $n \in \mathbb{N}^*$ a proof of $\forall X (X \Rightarrow X) \Rightarrow (X \Rightarrow X)$.

In Linear Logic, this yields a proof of $?(X^\perp \otimes X), !(X^\perp \otimes X)$, which is obtained from the sequent derivation of the previous subsection (in Figure 3.1) as follows:

$$\begin{array}{c} \vdots \\ \frac{\vdash X_1^\perp \otimes X_2, \dots, X_{m-5}^\perp \otimes X_{m-4}, X_{m-3}^\perp \otimes X_{m-2}, X_{m-1} \wp X_m^\perp}{\vdash ?(X_1^\perp \otimes X_2), \dots, X_{m-5}^\perp \otimes X_{m-4}, X_{m-3}^\perp \otimes X_{m-2}, X_{m-1} \wp X_m^\perp} ? \\ \vdots \\ \frac{\vdash ?(X_1^\perp \otimes X_2), \dots, ?(X_{m-5}^\perp \otimes X_{m-4}), X_{m-3}^\perp \otimes X_{m-2}, X_{m-1} \wp X_m^\perp}{\vdash ?(X_1^\perp \otimes X_2), \dots, ?(X_{m-5}^\perp \otimes X_{m-4}), ?(X_{m-3}^\perp \otimes X_{m-2}), X_{m-1} \wp X_m^\perp} ? \\ \frac{\vdash ?(X_1^\perp \otimes X_2), \dots, ?(X_{m-5}^\perp \otimes X_{m-4}), ?(X_{m-3}^\perp \otimes X_{m-2}), X_{m-1} \wp X_m^\perp}{\vdash ?(X_1^\perp \otimes X_2), \dots, ?(\mathbf{X}_1^\perp \otimes \mathbf{X}_2), X_{m-1} \wp X_m^\perp} \text{ctr.} \\ \vdots \\ \frac{\vdash ?(\mathbf{X}_1^\perp \otimes \mathbf{X}_2), X_{m-1} \wp X_m^\perp}{\vdash ?(\mathbf{X}_1^\perp \otimes \mathbf{X}_2), !(\mathbf{X}_3 \wp \mathbf{X}_4^\perp)} ! \end{array}$$

The literals in the conclusion are in bold and with a fresh labelling, to dwell on the fact that they do not stand only for themselves: $?(\mathbf{X}_1^\perp \otimes \mathbf{X}_2)$ is made of “strata” of contracted occurrences, and $!(\mathbf{X}_3 \wp \mathbf{X}_4^\perp)$ can be reproduced, it is the result of the computation. The Figure 3.3 is the introduction of the contractions in the proof net presented in Figure 3.2.

Informally, if we give n times 1 stick, we are able to obtain 1 pack of n sticks. What a low-level of computation ! But at least we are able —and this is the precise interest of Linear Logic— to handle resources within proofs.

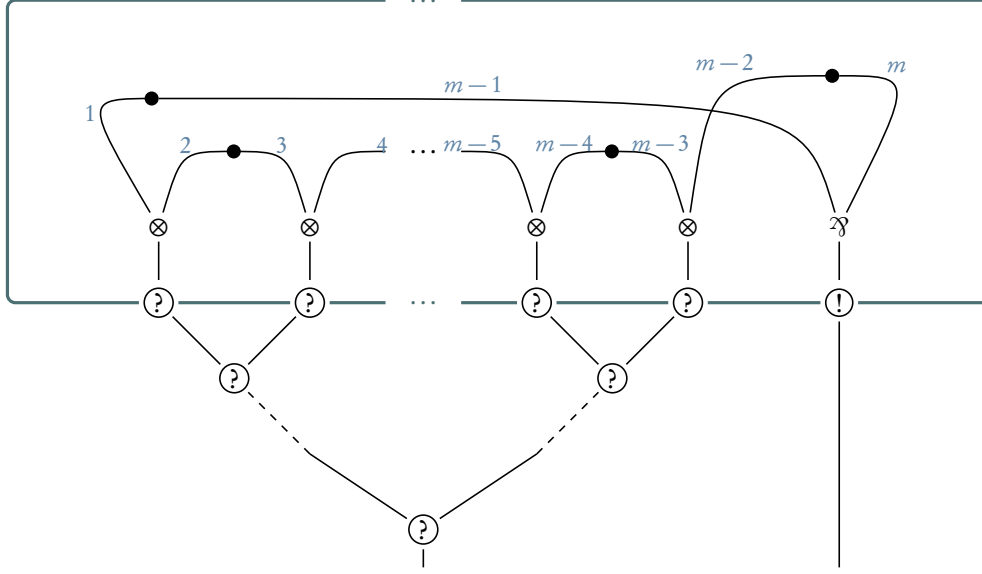
Contrary to the previous subsection, the type does not reflect all alone the integer encoded, it just reflects the mere fact that *an* integer has been encoded. More precisely, each tally stick $X_i^\perp \otimes X_j$ lives in a different *slice*, and all those slices are “compressed” into one single occurrence of $?(\mathbf{X}_1^\perp \otimes \mathbf{X}_2)$.

In this case, n occurrences of $X^\perp \otimes X$ are contracted, it should be more clear if we precise that the original proof of

$$X_1 \Rightarrow X_2, X_3 \Rightarrow X_4, \dots, X_{m-1} \Rightarrow X_{m-2} \vdash X_{m-1} \Rightarrow X_m$$

⁹As GIRARD likes to write, contraction is the “fingernail of infinity” [49, p. 12].

¹⁰The system F is the second-order λ -calculus, discovered independently by Jean-Yves GIRARD and John C. REYNOLDS, one may easily find a precise definition in textbooks [62, Chap. II].



This proof net permutes the applications of ! and ctr. compared to the derivation we gave. The ! rule needs all the occurrences in the context to be under a ? modality, hence the “box”, a classical tool we did not introduce.

Figure 3.3: Unaries as proof net, with contraction

can be renumbered as

$$X_{1,1} \Rightarrow X_{2,1}, X_{1,2} \Rightarrow X_{2,2}, \dots, X_{1,n} \Rightarrow X_{2,n} \vdash X_{3,0} \Rightarrow X_{4,0}$$

And thus we contract all the occurrences of $X_{1,y}$ into $\mathbf{X}_1$, and all the occurrences of $X_{2,y}$ into $\mathbf{X}_2$.

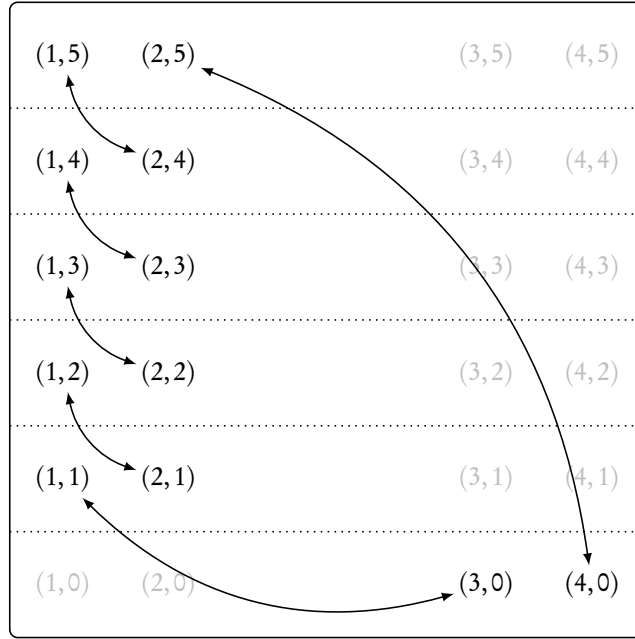
We are beginning to introduce two-dimensional considerations: we have 4 literals, and each of them contains $n + 1$ “floors”¹¹: the “ground floor”¹² is reserved for the result, and every tally stick lives on a different storey i.e. on a floor that is not the ground floor.

Note that nothing in the Logic prevents from displaying the first stick and the result on the same floor: we preferred to separate them, in order to clearly distinguish the beginning ($X_{3,0}$) and the end ($X_{4,0}$) of the integer represented.

We borrow from SEILLER’s thesis [I21, p. 251] a quite explanatory figure, representing the proof of the integer 5:

¹¹Although “slice” is the consecrated term for this notion, it is a bit misleading in this precise case. Contraction is handled as superimposition of copies of occurrences, as in the classical GoI framework, but that does not justify why (3,0) and (4,0) lives in different slices.

¹²To whom we give the number 0, as in the British convention.



It would need too much material to explain in details this figure, which actually represent a *thick graph* in the graph-theoretical representation introduced by SEILLER [121]. Two literals side-by-side are in black if they are “active” (present in the proof), and for each of them there is a $\Rightarrow$ -connective between them in the conclusion (before the contractions). The arrows represent the axiom links: remember for instance that $X_{3,0}$ and $X_{1,1}$ were respectively numbered with X_{m-1} and X_1 previously, and refer to Figure 3.1.

Remark that $X_{1,0}$, $X_{2,0}$, $X_{3,x}$ and $X_{4,x}$ for $0 < x \leq n$ are necessary empty, because no contracted occurrence will ever rest on those “slots”.

Now for the matricial representation, we still interpret the axiom rule as the symmetric exchange between two occurrences, but we split the information between several coefficient of M , a 4×4 matrix, as the number of occurrences of $\mathbf{X}$. Each of these coefficient is a $(n+1) \times (n+1)$ matrix, as the number of “floors” in the proof.

So far, we did not really took into account the case where $n = 0$, for it is at the same time quite different and not really interesting. Simply notice that 0 is quite an exception, for its proof is simply an axiom and a weakening, so there is only one symmetric exchange, and we represent the proof of 0 with

$$M_0 := \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{pmatrix}$$

The general case —i.e. $n \neq 0$ — is a bit more complex, let us remember that the axiom links are between

$$\begin{array}{ll}
X_{3,0} \text{ and } X_{1,1}, & \dots \\
X_{2,1} \text{ and } X_{1,2}, & X_{2,(n-1)} \text{ and } X_{1,n}, \\
X_{2,2} \text{ and } X_{1,3}, & X_{2,n} \text{ and } X_{4,0}.
\end{array}$$

We will represent any integer $n \in \mathbb{N}^*$ with the same canonical 4×4 matrix M :

$$M := \begin{pmatrix} 0 & v_n & u_n & 0 \\ v_n^t & 0 & 0 & w_n^t \\ u_n^t & 0 & 0 & 0 \\ 0 & w_n & 0 & 0 \end{pmatrix}$$

The matrix M is such that the entry $M[i, j]$ in the i th row and j th column encodes the exchanges between occurrences labelled with $X_{i,x}$ and $X_{j,y}$ for all $0 \leq x \leq n$ and $0 \leq y \leq n$. So for instance, $M[i, i]$ is necessarily empty, and $M = (M)^t$ because the exchanges are symmetric.

The coefficients of M — u_n , v_n and w_n — are $(n+1) \times (n+1)$ matrices that encodes¹³ respectively the exchanges between $X_{1,x}$ and $X_{3,x}$ —the beginning of the integer—, $X_{1,x}$ and $X_{2,x}$ —the body of the integer—, $X_{4,x}$ and $X_{2,x}$ —the end of the integer. Their size vary with n the integer represented, but they always have the same shape:

$$u_n := \begin{pmatrix} 0 & 0 & \dots & 0 & 0 \\ 1 & 0 & \dots & 0 & 0 \\ 0 & 0 & \dots & 0 & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & \dots & 0 & 0 \end{pmatrix} \quad v_n := \begin{pmatrix} 0 & 0 & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & \dots & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \dots & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & \dots & 0 & 1 & 0 \end{pmatrix} \quad w_n := \begin{pmatrix} 0 & \dots & 1 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix}$$

What is really interesting is their properties as matrices. We need before that to define the projections $\pi_{n,i}$ of dimension $(n+1) \times (n+1)$ with $0 \leq i \leq n$.

$$\pi_{n,i} := \begin{pmatrix} & & i & & \\ 0 & \dots & 0 & \dots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & 1 & \dots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & \dots & 0 \end{pmatrix} \quad i$$

Remark that for all i , $\pi_{n,i}^2 = \pi_{n,i}^t = \pi_{n,i}$ and that $\sum_{0 \leq j \leq n} \pi_{n,j} = \text{Id}_{n+1}$.

¹³We take as a convention that those matrice's rows and columns are numbered from 0 to n , whereas M was numbered from 1 to 4. This may seem illegitimate, but this help us to conserve the same numbering of the literals and floors.

Then one can remark that u_n , v_n and w_n are such that:

$$\begin{aligned} u_n u_n^t &= \pi_{n,1} \\ v_n v_n^t &= \pi_{n,2} + \dots + \pi_{n,n} \\ w_n w_n^t &= u_n^t u_n = \pi_{n,0} \\ w_n^t w_n &= \pi_{n,n} \end{aligned}$$

It will become of greatest importance when we will represent u_n , v_n and w_n as operators: instead of the mere transposition we will conserve those equalities when we use the conjugate-transposition $(\cdot)^*$ and that will establish that those three operators are partial isometries.

For the time being, and that will conclude this first part, let us remark that the following equalities hold:

$$\begin{aligned} u_n \pi_{n,i} &= \begin{cases} u_n & \text{if } i = 0 \\ 0 & \text{elsewhere} \end{cases} & \pi_{n,i} u_n &= \begin{cases} u_n & \text{if } i = n \\ 0 & \text{elsewhere} \end{cases} \\ v_n \pi_{n,0} &= v_n \pi_{n,n} = \pi_{n,0} v_n = \pi_{n,n} v_n = 0 \\ \pi_{n,n} v_n &= v_n \pi_{n,n-1} \\ v_n &= \pi_{n,2} v_n + \dots + \pi_{n,n} v_n = v_n \pi_{n,1} + \dots + v_n \pi_{n,n-1} \\ w_n \pi_{n,i} &= \begin{cases} w_n & \text{if } i = n \\ 0 & \text{elsewhere} \end{cases} & \pi_{n,i} w_n &= \begin{cases} w_n & \text{if } i = 0 \\ 0 & \text{elsewhere} \end{cases} \\ w_n &= u_n^t (v_n^t)^{-1} \\ \pi_{n,i+1} &= v_n \pi_{n,i} v_n^t \text{ for } 1 \leq i < n \end{aligned}$$

It is not of major interest for the time being to understand them precisely, they are just here to help the reader to grasp that we are able to go back and forth between projections and those coefficients: their sum “is a sort of circular permutation of $n+1$ projections $\pi_{0,n}, \dots, \pi_{n,n}$ ”, they “organise a ‘round-trip’” (GIRARD [59, pp. 253–254]).

They make it possible to cycle through the encoded representation. They are the core of the representation, for we represent the data by its dynamic, and only its dynamics has to be preserved by the embedding in the hyperfinite factor. We won’t develop that point here, all interesting equalities will be developed in Section 3.3 for the binary representation.

This representation, with the contraction, is a gain in term of uniformity,¹⁴ for every integer is represented with a 4×4 matrix, but if we look closely, those matrices are equivalent to $4(n+1) \times 4(n+1)$ matrices. To obtain a true uniformity, we will embed those matrices in the Π_1 hyperfinite factor $\mathfrak{R}$, as detailed in Proposition 3.4.1 for the binary representation: the operators will represent integers no matter their size, but conserve their properties.

We postpone this embedding for the case that interests us the most: the representation of the binary integers.

¹⁴Uniformity being understood in this case as “one object to represent them all”.

3.3 Binary Integers

We make one last time the construction from integers to matrices in this section, but it is to prepare the ground for Section 3.4, which pushes the construction one step further by embedding this representation in the hyperfinite factor.

In the previous chapter, we encoded binary words as a sequence of proof nets encoding bits. The information whereas the bit was 0 or 1 was encoded in the planarity of the proof net—in the presence or the absence of a permutation, or an exchange, between two literals.

Here we will proceed differently: we will encode binary words into a single proof net and the planarity will play no role. The information whereas the i th bit is a 0 or a 1 will rest in the order of the contractions. Every binary word of arbitrary size will have the same type $?(X \otimes X^\perp), ?(X \otimes X^\perp), !(X^\perp \otimes X)$. Whereas the conclusion is always the same, the number of application of the contraction rule and their principal formulae will differ from one proof to another.

We identify binary words with integers, using the “binary strings to integers” classical bijection:

For $a_k a_{k-1} \dots a_2 a_1$ a binary list of length k , $a_j \in \{0, 1\}$ for all $0 \leq j \leq k$,

$$a_k a_{k-1} \dots a_2 a_1 \rightarrow (a_k \times 2^{k-1}) + (a_{k-1} \times 2^{k-2}) + \dots + (a_2 \times 2^1) + (a_1 \times 2^0) = \sum_k (a_k \times 2^{k-1})$$

Here we won’t have a bijection, because for instance $0a_k \dots a_1$ will be accepted as a representing the same integer as $a_k \dots a_1$. Anyway, we can always uniquely associate an integer to every binary word, and we will sometimes commit the small abuse of speaking of *the* representation of an integer as a binary word.

The reader who read the previous section should not get confused: in the unary representation, the size of the representation *was* the integer represented, so we spoke of n as well as the integer and as its size. This setting is of course different, for several integers $n_1, \dots, n_{2^k}$ can all be represented with k bits. We will always assume in the following that the length of the binary word under study is $k \neq 0$.

Binaries as proofs

Following the Church-Howard representation, the binary word $a_k \dots a_1$ for $a_1, \dots, a_k \in \{0, 1\}^k$ is represented with the λ -term $\lambda f_0 \lambda f_1 \lambda x \cdot f_{a_k} (f_{a_{k-1}} (\dots (f_{a_1} x) \dots))$. Given two function f_0 —a.k.a. “append a 0 to the list”— and f_1 —a.k.a. “append a 1 to the list”—, the binary word is encoded in the order of application of those two functions. Of course, we need only one copy of each of those functions: in Linear Logic, this will correspond to contractions.¹⁵

In ELL,¹⁶ binary lists are typed with

$$\forall X \ ! (X \multimap X) \multimap (! (X \multimap X) \multimap ! (X \multimap X))$$

To a binary integer corresponds a proof of the —annotated— Linear Logic sequent

$$\forall X \left(?(X(0_i) \otimes X(0_o)^\perp) \wp \left(?(X(1_i) \otimes X(1_o)^\perp) \wp ! (X(S)^\perp \wp X(E)) \right) \right) \quad (3.1)$$

¹⁵Yes, to contraction, for the proof is built upside-down and is one-sided: so what could normally look like a duplication —“gives me one function and I will apply it several times”— is in reality a contraction—“I used the function several times but it will appear only once in the conclusion”.

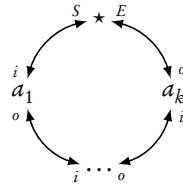
¹⁶This notion and all the references were introduced p. 18.

We should take some time to explain those annotations: $X(0_i) \otimes X(0_o)$ plays the role of adding a 0 at the end of the list, and it has an *input* (the original list) and an *output* (the list with a 0 appended), respectively at the left and the right of the $\otimes$ -connective. The other “function” $X(1_i) \otimes X(1_o)$ plays the same role for “append a 1 to the list”.

The last part of this formula $-X(S)^\perp \wp X(E)-$ is the “exterior interface”: it has also an output (the “ S ” part, as in *start*) and an input (the “ E ” part, as in *end*).¹⁷ It acts like a marker of the beginning *and* the end of the word. We use the symbol $\star$ to represent this marker, from which we may *in one step* go to the end or to the beginning of the word.

Each element is connected by its output node to its successor’s input node and by its input node to its predecessor’s output node. The successor of the last element and the predecessor of the first element is $\star$.

So a binary word will be circular, we represent this fact by using the notation $\star a_k \dots a_1 \star$. The bits are numbered in the reverse order: the leftmost bit is the last (the end of the word, the k th one) and the rightmost bit is the first (the start of the word). By convention, the $k + 1$ th and the 0th bits are $\star$. Or, more graphically, and with labels regarding the input / output structure:



One could see above that S can be read as $\star$ ’s output and E as $\star$ ’s input.

The size and the structure of the derivation of such a list change a lot in some cases: the empty list uses the weakening rule twice, and the proof representing a list containing only 0s (resp. only 1s) uses a weakening to introduce $\wp(X(1i) \otimes X(1o)^\perp)$ (resp. $\wp(X(0i) \otimes X(0o)^\perp)$). The reader may refer to earlier works [I21, I0, p. 254–257] to see the proofs and the associated graphs in those cases.

We won’t take them into account, for we may safely assume that we are always working with binary words containing 0 and 1. Any sequence $11 \dots 11$ can safely be replaced by $011 \dots 11$, and we can always suppose we don’t need any representation of 0.

The derivation of the sequent (3.1) is given in all generality in Figure 3.4. We labelled the variables in order to distinguish between the different occurrences of the variable X . This is necessary because we need to keep track of which formulae are principal in the contraction rule. Without this information, a lot of proofs could not be distinguished.

Example 3.3.1 (Eleven (1/3)). We propose in Figure 3.5 an example of a binary word $\star 001011 \star$ represented as a proof net and as a circular word. This sequence—which is one of the representation of eleven—will be a running example in this first part of this chapter and should help the reader to grasp the general shape of the construction.

There is one last representation that we chose not to develop: the *interaction graphs* invented by SEILLER [I21, I20, I19]. This approach is enlightening as it allows us to see the Geometry of Interaction as a combinatorial object. Yet we do not develop this setting here, for it would need to introduce too many objects at the same time. A first taste of this representation was given page 65: we refer the reader

¹⁷“Start” and “End” rather than the more classical “Head” and “Tail”, mainly because in this setting, the head *is* the tail, and that could be misleading.

We give below the derivation of the proof corresponding to $\star a_k a_{k-1} \dots a_1 \star$, which could be read backward and with the labels as $S a_{i_1} a_{o_1} \dots a_{i(k-1)} a_{o(k-1)} \dots a_{i_k} a_{o_k} E$.

$$\begin{array}{c}
\frac{\frac{\frac{}{\vdash X(S)^\perp, X(a_{i_1})} \text{ax.} \quad \frac{}{\vdash X(a_{o_1})^\perp, X(a_{i_2})} \text{ax.}}{\vdash X(S)^\perp, X(a_{i_1}) \otimes X(a_{o_1})^\perp, X(a_{i_2})} \otimes \quad \frac{}{\vdash X(a_{o_2})^\perp, X(a_{i_3})} \text{ax.}}{\vdash X(S)^\perp, X(a_{i_1}) \otimes X(a_{o_1})^\perp, X(a_{i_2}) \otimes X(a_{o_2})^\perp, X(a_{i_3})} \otimes \\
\vdots \\
\frac{\frac{\vdash X(S)^\perp, X(a_{i_1}) \otimes X(a_{o_1})^\perp, \dots, X(a_{i(k-1)}) \otimes X(a_{o(k-1)})^\perp, X(a_{i_k})}{\vdash X(S)^\perp, X(a_{i_1}) \otimes X(a_{o_1})^\perp, \dots, X(a_{i(k-1)}) \otimes X(a_{o(k-1)})^\perp, X(a_{i_k}) \otimes X(a_{o_k})^\perp, X(E)} \otimes \quad \frac{}{\vdash X(a_{o_k})^\perp, X(E)} \text{ax.}}{\vdash X(a_{i_1}) \otimes X(a_{o_1})^\perp, \dots, X(a_{i(k-1)}) \otimes X(a_{o(k-1)})^\perp, X(a_{i_k}) \otimes X(a_{o_k})^\perp, X(S)^\perp \wp X(E)} \wp \\
\frac{}{\vdash \wp(X(a_{i_1}) \otimes X(a_{o_1})^\perp), \dots, \wp(X(a_{i(k-1)}) \otimes X(a_{o(k-1)})^\perp), \wp(X(a_{i_k}) \otimes X(a_{o_k})^\perp), !(X(S)^\perp \wp X(E))} ! \\
\vdots \\
\frac{\frac{\vdash \wp(X(0_i) \otimes X(0_o)^\perp), \wp(X(1_i) \otimes X(1_o)^\perp), !(X(S)^\perp \wp X(E))}{\vdash \wp(X(0_i) \otimes X(0_o)^\perp), (\wp(X(1_i) \otimes X(1_o)^\perp) \wp !(X(S)^\perp \wp X(E)))} \wp \\
\frac{\vdash \wp(X(0_i) \otimes X(0_o)^\perp) \wp (\wp(X(1_i) \otimes X(1_o)^\perp) \wp !(X(S)^\perp \wp X(E)))}{\vdash \forall X (\wp(X(0_i) \otimes X(0_o)^\perp) \wp (\wp(X(1_i) \otimes X(1_o)^\perp) \wp !(X(S)^\perp \wp X(E))))} \forall
\end{array}$$

The formulae $\wp(X(a_{i_c}) \otimes X(a_{o_c})^\perp)$ and $\wp(X(a_{i_d}) \otimes X(a_{o_d})^\perp)$ are contracted iff $a_c = a_d$. After the applications of the contraction rule, we have only two occurrences of $\wp(X(a_i) \otimes X(a_o)^\perp)$ for $a \in \{0, 1\}$, gathering the occurrences of 0 and the occurrences of 1.

Figure 3.4: Binary integers as derivations

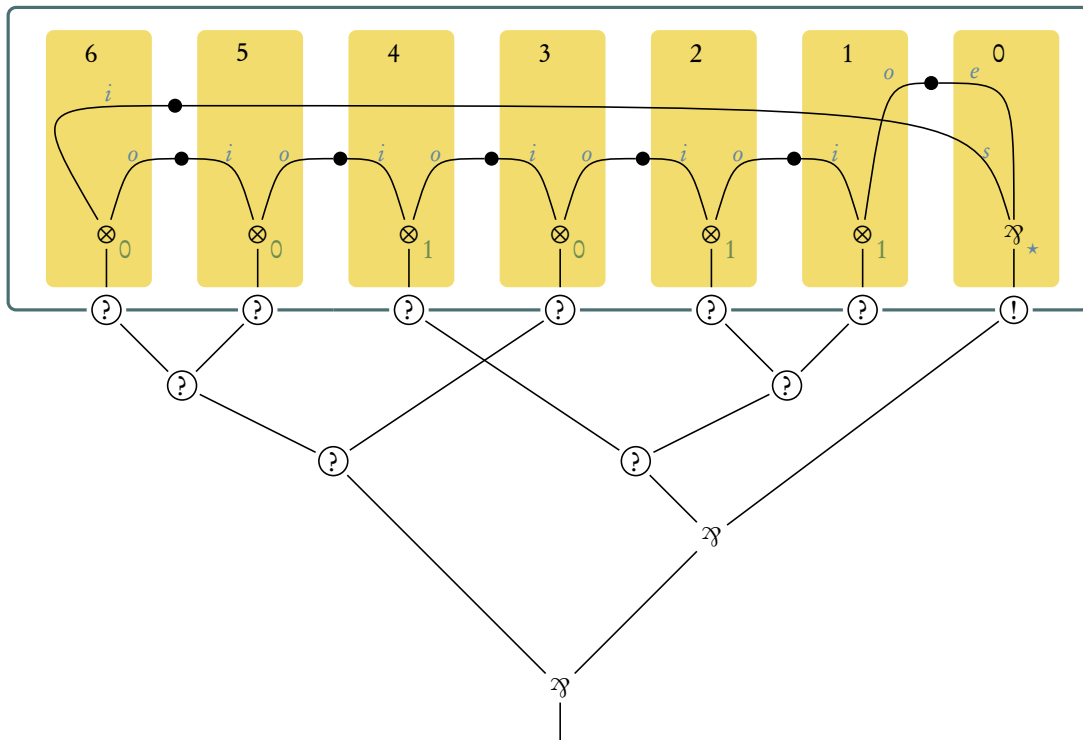
to this subsection or to an earlier version of this work [10] for a better insight on binary words as interaction graphs.

In the forthcoming setting, proof nets are not our main objects, we will use matrices as representations of (proofs of) integers. Yet, it would probably be wiser to avoid the non-planarity of our proof nets, since in this case it only depends on a convention regarding where is the input and output of our word. That would require a careful renaming of the variable corresponding to S and E in this work, but would give another property to our representation and probably make it gain in clarity. We did not make this small change —suggested by Paul-André MELIÈS— since our next representation is not sensitive to that property of proof nets.

Binaries as matrices

The idea to represent proofs as matrices comes from the simple remark that all cut-free proofs in LL of same conclusion differs essentially only on their axioms. The reader who looked at the previous section should be used to this idea, but it is not mandatory to understand this motivation to grasp the following construction.

Any binary word will be represented by the same 6×6 matrix where one row and one column is dedicated to a literal of Equation 3.1. We recall that those literals handle the information about the input and output structure of 0, 1 and $\star$. So for instance the coefficient at the intersection of the “input of 0” row and the “output of 1” will contains the positions in the sequence where a 1 is followed by a 0.



We indicate the “binary values” (0 and 1) next to the logical connectives $\otimes$, but their values comes from the order of the contractions.

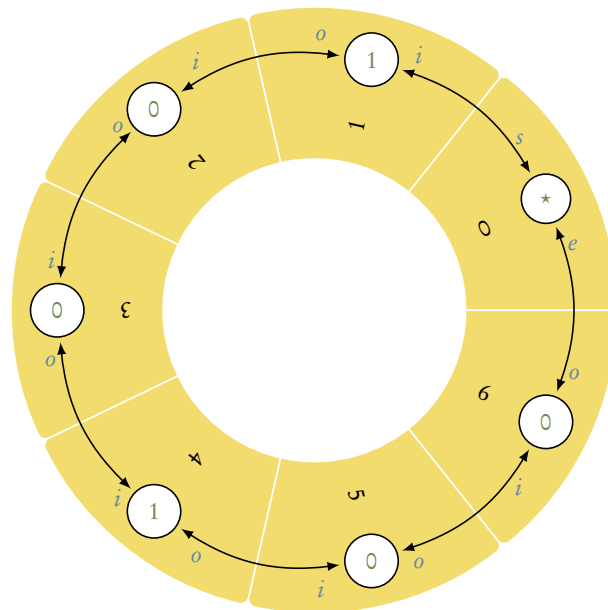


Figure 3.5: The binary sequence $*001011*$ as a proof net and as a circular string

Definition 3.3.1 (Matricial representation of a list). Given a binary representation $\star a_k \dots a_1 \star$ of an integer n , it is represented by M_n , a 6×6 block matrix¹⁸ of the following form:

$$M_n = \begin{pmatrix} \overbrace{0 \quad l_{00} \quad 0 \quad l_{10} \quad l_{S0} \quad 0}^0 & \overbrace{l_{00}^t \quad 0 \quad l_{01}^t \quad 0 \quad 0 \quad l_{0E}^t}^1 & \overbrace{l_{10}^t \quad 0 \quad l_{11}^t \quad 0 \quad 0 \quad l_{1E}^t}^\star \\ \overbrace{0 \quad l_{01} \quad 0 \quad l_{11} \quad l_{S1} \quad 0}^0 & \overbrace{l_{10} \quad 0 \quad l_{11} \quad 0 \quad 0 \quad 0}^1 & \overbrace{l_{S0} \quad 0 \quad l_{S1} \quad 0 \quad 0 \quad 0}^\star \\ \overbrace{0 \quad l_{0E} \quad 0 \quad l_{1E} \quad 0 \quad 0}^\star \end{pmatrix}$$

whose entries are $(k+1) \times (k+1)$ matrices¹⁹ defined —for $u, v \in \{0, 1\}$ — by:

- $l_{uv}[i, i+1] = 1$ iff $a_i = u$ and $a_{i+1} = v$, $l_{uv}[a, b] = 0$ otherwise;
- $l_{uE}[0, 1] = 1$ iff $a_k = u$, and $l_{uE}[a, b] = 0$ otherwise,
- $l_{Sv}[k, 0] = 1$ iff $a_1 = v$, and $l_{Sv}[a, b] = 0$ otherwise.

Remark that —as the output of a bit is always connected to the input of its direct neighbor— $l_{uv}[a, b] = 1$ implies that $b = a + 1$. This was for the body of the word. Regarding the start (resp. the end) of the word, only one among l_{S0} , l_{S1} (resp. l_{0E} , l_{1E}) is different from 0: only one value among 0, 1 may be connected to S (resp. E), that is, the sequence necessarily starts (resp. ends) with a 0 or with a 1.

A sketch in Figure 3.6 should help the reader to grasp some properties of this matrix, as well as its inner structure.

Example 3.3.2 (Eleven (2/3)). Our example $\star 001011 \star$ is represented by M_{11} whose coefficients²⁰ are as follows:

- $l_{00}[a, b] = 0$ except for $l_{00}[5, 6] = 1$
- $l_{01}[a, b] = 0$ except for $l_{01}[3, 4] = 1$
- $l_{10}[a, b] = 0$ except for $l_{10}[2, 3] = l_{10}[4, 5] = 1$
- $l_{11}[a, b] = 0$ except for $l_{11}[1, 2] = 1$
- $l_{1E}[0, 1] = 1$, $l_{S0}[6, 0] = 1$, and all the other coefficients of l_{1E} and l_{S0} are 0
- $l_{0E} = l_{S1} = 0$

This representation of binary integers is still “non-uniform”: the size of the matrix — $6(k+1) \times 6(k+1)$ — depends on the size k of the binary string representing the integer.

To get a uniform representation of integers, we still need to embed the coefficients of M_n in $\mathfrak{R}$. But before that, we should dwell on some properties of this matricial representation.

¹⁸The embraces are only pedagogical.

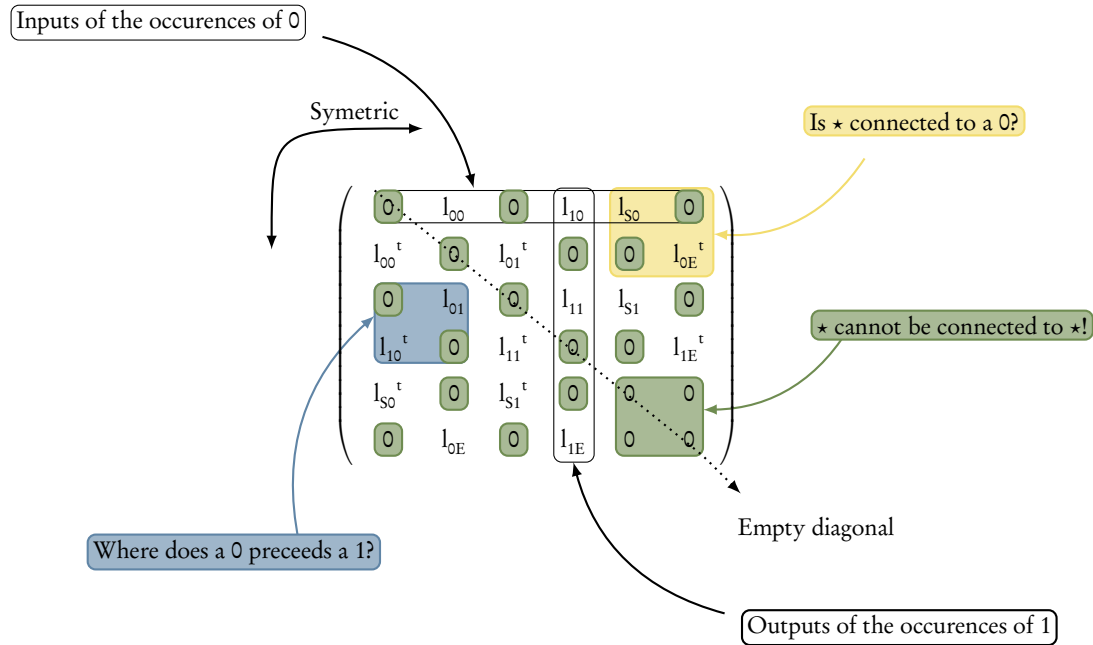
¹⁹Matrices whose rows and columns are numbered from 0 to k , for convenience.

²⁰Remember that the first element of the word is the rightmost one.

The semantics of the coefficients of M_n is, given a string where a_k is the last bit, a_1 the first, a and b any two bits:

l_{ab}	Goes from the output of $a \in \{0, 1\}$ to the input of $b \in \{0, 1\}$.
l_{ab}^t	Goes from the input of $b \in \{0, 1\}$ to the output of $a \in \{0, 1\}$.
l_{sa_k}	Goes from the output of $\star (= S)$ to the input of $a_k \in \{0, 1\}$.
$l_{sa_k}^t$	Goes from the input of $a_k \in \{0, 1\}$ to the output of $\star (= S)$.
$l_{a_1 E}$	Goes from the output of $a_1 \in \{0, 1\}$ to the input of $\star (= E)$.
$l_{a_1 E}^t$	Goes from the input of $\star (= E)$ to the output of $a_1 \in \{0, 1\}$.

Each of those coefficient is a $(k+1) \times (k+1)$ matrix: the dimension corresponds to the size of the string encoded plus an additional room to handle $\star$. The matrix M_n is a 6×6 matrix—as the number of occurrence of literals in the conclusion of the proof—organized as follows:



The values **in a rectangle** are 0 because an input (resp. output) will never be connected to another input (resp. output), or because $\star$ will never be connected to itself.

The reader should of course refer to [Definition 3.3.1](#) for a precise definition of M_n the matricial representation of n .

Figure 3.6: Some explanations about the matricial representation of a binary word

Some remarkable properties of the matricial representation

For $k \in \mathbb{N}^*$ and $0 \leq i \leq k$, we define $\pi_{k,i}$ to be the $(k+1) \times (k+1)$ matrix such that

$$\pi_{k,i}[a, b] = \begin{cases} 1 & \text{if } a = b = i \\ 0 & \text{elsewhere} \end{cases} \quad \text{i.e.} \quad \pi_{k,i} := \begin{pmatrix} & & i & & \\ 0 & \dots & 0 & \dots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & 1 & \dots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & \dots & 0 \end{pmatrix}_i$$

We will in the following omit the subscript k , for the dimension of the projection should always be clear from context. Those projections shares a lot with the encoding we defined, as the small list of equalities that follow shows. This subsection can be seen as a small break in our journey, where we look at how the objects we defined behave.

$$\begin{aligned} \pi_1 &= (l_{0E} + l_{1E})(l_{0E} + l_{1E})^t \\ \pi_0 &= (l_{0E} + l_{1E})^t(l_{0E} + l_{1E}) \end{aligned}$$

Depending on the last bit, l_{0E} or l_{1E} is

$$\begin{pmatrix} 0 & 0 & \dots & 0 \\ 1 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix}$$

The other being the 0 matrix. And it is quite clear that

$$\begin{aligned} \begin{pmatrix} 0 & 0 & \dots & 0 \\ 1 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} \begin{pmatrix} 0 & 1 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} &= \begin{pmatrix} 0 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} \\ \begin{pmatrix} 0 & 1 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 & \dots & 0 \\ 1 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} &= \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 \end{pmatrix} \end{aligned}$$

We have a similar result regarding the beginning of the word represented:

$$\begin{aligned} \pi_0 &= (l_{S0} + l_{S1})(l_{S0} + l_{S1})^t \\ \pi_k &= (l_{S0} + l_{S1})^t(l_{S0} + l_{S1}) \end{aligned}$$

According to the value of the first bit, one among l_{S0} and l_{S1} is the 0 matrix, the other being

$$\begin{pmatrix} 0 & \dots & 1 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix}$$

And a simple computation proves our claim.

$$\begin{pmatrix} 0 & \dots & 1 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix} \begin{pmatrix} 0 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 1 & \dots & 0 \end{pmatrix} = \begin{pmatrix} 1 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix}$$

$$\begin{pmatrix} 0 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 1 & \dots & 0 \end{pmatrix} \begin{pmatrix} 0 & \dots & 1 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 0 \end{pmatrix} = \begin{pmatrix} 0 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & 1 \end{pmatrix}$$

Informally, π_0 is the bridge between the start and the end of our word. To read the beginning of the word, we should go toward π_k , whereas π_1 gives information about the end of the sequence.

The representation of the body of the word has also a tight link with the projections. First, we should remark that the sum of the matrices has always the same structure, in fact it is the same that the one used to represent unary structures.

$$(l_{S0} + l_{S1}) + (l_{00} + l_{01} + l_{10} + l_{11}) + (l_{1E} + l_{0E}) = \begin{pmatrix} 0 & 0 & 0 & \dots & 0 & 0 & 1 \\ 1 & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & 1 & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & 1 & \dots & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 1 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 1 & 0 \end{pmatrix}$$

This is not surprising, for it simply tells us that every single bit “is defined”: for every one of them, there is one and only one matrix that contains information about it. The part regarding the body of the word

$$(l_{00} + l_{01} + l_{10} + l_{11}) = \begin{pmatrix} 0 & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 \\ 0 & 1 & 0 & \dots & 0 & 0 & 0 \\ 0 & 0 & 1 & \dots & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 1 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 1 & 0 \end{pmatrix}$$

has an internal structure that links it to a sum of projections:

$$(l_{00} + l_{01} + l_{10} + l_{11})(l_{00} + l_{01} + l_{10} + l_{11})^t = \pi_2 + \pi_3 + \dots + \pi_k$$

$$(l_{00} + l_{01} + l_{10} + l_{11})^t(l_{00} + l_{01} + l_{10} + l_{11}) = \pi_1 + \pi_2 + \dots + \pi_{k-1}$$

Moreover, for $k > i > 0$,

$$\pi_{i+1} = (l_{00} + l_{01} + l_{10} + l_{11})\pi_i(l_{00} + l_{01} + l_{10} + l_{11})^t \quad (3.2)$$

This equality tells us somehow how to go read the next bit. Of course, this equality does not apply for the first or the last bit, because those matrices contains no information about those bits. This is

expressed by $(l_{00} + l_{01} + l_{10} + l_{11})^k = 0$. This can be easily checked:

$$(l_{00} + l_{01} + l_{10} + l_{11})^2 = \begin{pmatrix} 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & \dots & 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 1 & 0 & 0 \end{pmatrix}$$

and so

$$(l_{00} + l_{01} + l_{10} + l_{11})^{k-1} = \begin{pmatrix} 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & \dots & 0 & 0 & 0 & 0 \end{pmatrix}$$

This dynamics is interesting not only in itself, but also because we will use it as a *definition* in the next section: what represents an integer is not its size or the shape of those matrices, but their *dynamic* when iterated or multiplied with projections, something that can be perfectly expressed by operators.

3.4 How to Compute

We will first explain how it is possible to represent integers in the hyperfinite factor, relying on the classical embedding of the matrix algebra into the type II_1 hyperfinite factor $\mathfrak{R}$. As you may guess, this embedding will strongly rely on the previous construction. A first approach on how matrices could be avoided by using C^* algebra was already developed [50, Proposition II, p. 58].

We will then spend some time discussing the representation of computation (or more precisely, acceptance) and then explain how programs can be represented in this setting. This part gives the answers to our concerns regarding *uniformity* and *objectivity*: we want a single kind of object to represent all the integers, and at the time we want the program to interact in the same way no matter the representative chosen.

Binaries as operators

We will embed the $(k+1) \times (k+1)$ matrices of Definition 3.3.1 in the hyperfinite factor $\mathfrak{R}$ in order to have a uniform representation of the integers: an integer —no matter its size— will be represented by an operator in $M_6(\mathfrak{R})$ fulfilling some properties. To express them we define, given a sequence $\star a_k \dots a_1 \star$

representing an integer n and for $j, l \in \{0, 1\}$, the sets:

$$I_{jl}^n = \{m \mid a_m = j, a_{m+1} = l, 1 \leq m \leq k\}$$

$$I_{Sj}^n = \{1\} \text{ if } a_1 = j, \emptyset \text{ elsewhere}$$

$$I_{jE}^n = \{k\} \text{ if } a_k = j, \emptyset \text{ elsewhere}$$

Roughly speaking, I_{Sj}^n (resp. I_{jE}^n) tells us about the first (resp. last) bit of the sequence and I_{jl}^n is the set of sequences of a j followed by a l . We let $\blacklozenge \in \{00, 01, 10, 11, S0, S1, 0E, 1E\}$ as a shortcut.

Those sets allow us to be more precise regarding the formulation of Equation 3.2:

$$\sum_{d \in I_{\blacklozenge}} l_{\blacklozenge} \pi_d = l_{\blacklozenge} \quad (3.3)$$

$$\sum_{d \in I_{\blacklozenge}} \pi_{d+1} l_{\blacklozenge} = l_{\blacklozenge} \quad (3.4)$$

This explain in hard-to-read terms a very simple fact that our example will help to grasp.

Example 3.4.1 (Eleven (3/3)). Our sequence $\star 001011\star$ yields

$$I_{S0} = \emptyset \quad I_{S1} = \{1\} \quad I_{0E} = \{6\} \quad I_{1E} = \emptyset \quad I_{00} = \{5\} \quad I_{01} = \{3\} \quad I_{11} = \{1\} \quad I_{10} = \{2, 4\}$$

Remember (Example 3.3.2) for instance that in this case,

$$l_{10} = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

This comes from the fact that in the 2nd and the 4th positions (counting from right to left) of our sequence there is a 1 followed by a 0. And so Equation 3.3, when instantiated to $I_{\blacklozenge} = I_{10}$, gives

$$\begin{aligned} \sum_{d \in I_{10}} l_{10} \pi_d &= \sum_{i=2,4} l_{10} \pi_i \\ &= (l_{10} \pi_2) + (l_{10} \pi_4) \\ &= l_{10} \end{aligned}$$

This is simple to check, as

$$l_{10} \pi_2 = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

and as

$$l_{10}\pi_4 = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix}$$

The “converse” Equation 3.4 can easily be checked the same way.

Now for the representation by operators: we will keep the same notation, simply change the family (from roman to *italic*), i.e. the operator representing l_{10} will be denoted by l_{10} . The only change will be that instead of taking the mere transposed of a matrix $(\cdot)^t$, we will use the the conjugate-transpose of operators, denoted by $(\cdot)^*$. The Proposition 3.4.1 will justify this identical notation.

Definition 3.4.1 (Binary representation of integers). An operator $N_n \in M_6(\mathfrak{A})$ is a *binary representation* of an integer n if there exists projections $\pi_0, \pi_1, \dots, \pi_k$ in $\mathfrak{A}$ that satisfy $\sum_{i=0}^k \pi_i = 1$ and

$$N_n = \begin{pmatrix} 0 & l_{00} & 0 & l_{10} & l_{S0} & 0 \\ l_{00}^* & 0 & l_{01}^* & 0 & 0 & l_{0E}^* \\ 0 & l_{01} & 0 & l_{11} & l_{S1} & 0 \\ l_{10}^* & 0 & l_{11}^* & 0 & 0 & l_{1E}^* \\ l_{S0}^* & 0 & l_{S1}^* & 0 & 0 & 0 \\ 0 & l_{0E} & 0 & l_{1E} & 0 & 0 \end{pmatrix}$$

The coefficients of N_n are partial isometries fulfilling the following equations (where $\pi_{k+1} = \pi_0$).

$$l_{\diamond} = \sum_{i \in I_n^{\diamond}} \pi_{i+1} l_{\diamond} \pi_i \quad (\diamond \in \{00, 01, 10, 11, S0, S1, 0E, 1E\}) \quad (3.5)$$

$$\pi_0 = (l_{S0} + l_{S1})(l_{00} + l_{01} + l_{10} + l_{11})^{k-1}(l_{0E} + l_{1E}) \quad (3.6)$$

Equation 3.5 and Equation 3.6 are the core of the computational behaviour we expect, they sum up the equations of the previous subsection and *define* our representation of integers as operators.

For all $i \in \{1, \dots, k\}$, let $u_i = ((l_{00} + l_{01} + l_{10} + l_{11})^i (l_{0E} + l_{1E}))$. We should remark that $u_i u_i^* = \pi_i$ and that $u_i^* u_i = \pi_0$, i.e. that π_0 is equivalent²¹ to π_i .

Proposition 3.4.1 (Binary and matricial representations). *Given $N_n \in M_6(\mathfrak{A})$ a binary representation of the integer n , there exists an embedding $\theta : M_{k+1}(\mathbb{C}) \rightarrow \mathfrak{A}$ such that²² $\text{Id} \otimes \theta(M_n) = N_n$, where M_n is the matricial representation of n .*

The global picture of the embedding we are going to define is as follows:

²¹In the sense of the Murray-von Neumann equivalence (remember Definition A.2.3).

²²We denote by Id the identity matrix of $M_k(\mathbb{C})$. We will allow ourselves the same abuse of notation in the following statements and proofs in order to simplify the formulae.

$$\begin{array}{ccc} M_6 (M_{k+1}(\mathbb{C})) & & \\ \downarrow \text{Id} \otimes \theta & & \downarrow \theta \\ M_6 (\mathfrak{R}) & & \end{array}$$

It uses the fact that $M_{6(k+1)}(\mathbb{C})$ is isomorphic to $M_6(\mathbb{C}) \otimes M_{k+1}(\mathbb{C})$. We explained the main steps to embed a matrix algebra into the hyperfinite factor in [Definition A.2.7](#), and perform this construction anew below.

Proof. Let $N_n \in M_6(\mathfrak{R})$ be a binary representation of $n \in \mathbb{N}$, and $\pi_0, \dots, \pi_k$ the associated projections. We now define an embedding $\theta : M_{k+1}(\mathbb{C}) \rightarrow \mathfrak{R}$:

$$\theta : M \mapsto \sum_{i=0}^k \sum_{j=0}^k M[i, j] u_{i,j}$$

with:

$$u_{i,j} = \begin{cases} (l_{00} + l_{01} + l_{10} + l_{11})^{j-1} (l_{0E} + l_{1E}) & \text{if } i = 0 \\ (l_{00} + l_{01} + l_{10} + l_{11})^{j-1} & \text{if } i < j \text{ and } i \neq 0 \\ ((l_{00} + l_{01} + l_{10} + l_{11})^{i-1} (l_{0E} + l_{1E}))^* & \text{if } j = 0 \\ ((l_{00} + l_{01} + l_{10} + l_{11})^{i-1})^* & \text{if } i > j \text{ and } j \neq 0 \\ \pi_k & \text{if } i = j = k \end{cases}$$

We can easily check that the image by $\text{Id} \otimes \theta$ of the matrix M_n representing n is equal to N_n .

For instance, let $\diamond \in \{00, 01, 10, 11\}$,

$$\begin{aligned} \theta(l_\diamond) &= \sum_{i=0}^k \sum_{j=0}^k l_\diamond[i, j] u_{i,j} && \text{(By definition of } \theta) \\ &= \sum_{p \in I_\diamond} l_\diamond[p, p+1] u_{p,p+1} && \text{(By definition of } l_\diamond) \\ &= \sum_{p \in I_\diamond} l_\diamond[p, p+1] (l_{00} + l_{01} + l_{10} + l_{11})^{p-1} && \text{(By applying case 3.7b)} \end{aligned}$$

When we embed the transposed of $(l_{00} + l_{01} + l_{10} + l_{11})$, the computation is the same except that we use case 3.7d.

Regarding $\theta(l_{1E} + l_{0E})$, remember that $(l_{1E} + l_{0E})[i, j] = 0$ except if $i = 0$ and $j = 1$, so by applying case 3.7a, $\theta(l_{1E} + l_{0E}) = (l_{00} + l_{01} + l_{10} + l_{11})^0 (l_{0E} + l_{1E}) = (l_{0E} + l_{1E})$. The embedding of $(l_{1S} + l_{0S})$ is quite similar, by applying case 3.7c.

We have of course that $\theta(M^t) = (\theta(M))^*$, and $u_{i,j}$ is defined so that $u_{i,j}^* = u_{j,i}$. Remark also that θ is not always defined, for instance $\theta(u_{n,n})$ where $n \neq 0, k$ is not defined. **Q.E.D.**

Rather than having $(k+1) \times (k+1)$ matrices in a 6×6 matrix, we have partial isometries in a 6×6 matrix: partial isometries adapts themselves to the size of the word we represent. So having the same matrix of operators is a gain in terms of uniformity, as all the integers are represented by matrices of the same size. But at the same time, as any embedding $\theta : M_{k+1}(\mathbb{C}) \rightarrow \mathfrak{R}$ defines a representation of the integers, we have to check that they all are equivalent.

Proposition 3.4.2 (Equivalence of binary representations). *Given N_n and N'_n two binary representations of $n \in \mathbb{N}^*$, there exists a unitary $u \in \mathfrak{K}$ such that $(\text{Id} \otimes u)N_n(\text{Id} \otimes u)^* = N'_n$.*

Proof. Let $\pi_0, \dots, \pi_k$ (resp. $\nu_0, \dots, \nu_k$) be the projections and $l_\blacklozenge$ (resp. $l'_\blacklozenge$) the partial isometries associated to N_n (resp. N'_n). It is straightforward that π_0 and ν_0 are equivalent according to Murray-von Neumann definition, so there exists a partial isometry v such that $vv^* = \nu_0$ and $v^*v = \pi_0$. For all $0 \leq i \leq k$ we define the partial isometries:

$$v_i = \underbrace{((l'_{00} + l'_{01} + l'_{10} + l'_{11})^{i-1}(l'_{0E} + l'_{1E}))}_{\nu_0} v \underbrace{((l_{00} + l_{01} + l_{10} + l_{11})^{i-1}(l_{0E} + l_{1E}))^*}_{\pi_0^*}$$

We can easily check that:

$$v_i^* v_i = \pi_i \qquad v_i v_i^* = \nu_i$$

It follows that the sum $u = \sum_{i=0}^n v_i$ is a unitary and $(\text{Id} \otimes u)N_n(\text{Id} \otimes u)^* = N'_n$.

Q.E.D.

All the representations of the same integer are equivalent, from a mathematical point of view, but if we want to compute with those integers, we have to be sure that the algorithm will make no difference between several representations. This will be the subject of the third subsection of this section, but before that we take some time to study *how* integers and programs interact, because that interaction experienced several definitions.

Where programs and integers interact

This small subsection is here to justify some choice, to dwell on the advances we proposed and to link this representation to an other choice made regarding GoI. The equivalence of all the representations has some philosophical motivations: GIRARD wanted all the representations to be “egal”, none of them should be more “normal” than any other.²³ This notion of a norm as an unjustified constraint explains the title “Normativity in Logic”.

But of course, what makes two representations equal or not is their behaviour from a computational point of view. We want to define “observations”, programs acting on them, and we want them to be “objective”, i.e. insensitive to the representation chosen. In this context, “abnormality” should become relative to the evaluating context.

Observations and integers will be the same kind of objects, i.e. operators in the II_1 hyperfinite factor, and they will interact through plain multiplication. The question is: “How can we represent the end of the computation?” We will define an accept / reject pattern, the observations do not produce any complex output. This first criteria proposed was “When the determinant of the product is 0”, the answer we developed with SEILLER is “When the product gets to 0”, i.e. if it is nilpotent. To be fair, this second criteria was already in use in “Normativity in Logic” [59], but only in background.

Technically speaking, there was a first attempt to define a frame where programs and integers could nicely interact, using a commutation condition, but it was way too constraining. In this setting, an algorithm could express only trivial computation [59, Proposition 12.1].

²³In fact, there is a detail that goes again the success of this plan: there is still a norm on the representations. If we look closely to Proposition 3.4.2, we should see that all the representations of an integer n has to have the same number of bits. This contradicts lightly the introductory remark of Section 3.3, but does not really alter the whole construction.

The interaction used by GIRARD was based on FUGLEDE-KADISON determinant, a generalization of the usual determinant of matrices that can be defined in a type II_1 factor. This requires to define “normalised” determinant and trace, for them to be independent on the dimension of the matrix under treatment. In Proposition 3.4.2, we showed that two representations were unitary equivalent, as the trace of a unitary is 1, that let the determinant of the representation of the integer unchanged. This used to be crucial when cancellation of determinant was used as a marker for the end of the computation. It was also interesting for the trace of the projections are independent of their size, and those “normalised” traces and determinants are invariant under the embedding of Proposition 3.4.1.

Those notions are nicely developed in SEILLER’s thesis [121, chap. 4.2] and it was proven, thanks notably to a technical lemma [120, Lemma 61], that this condition on determinant was equivalent to the condition on nilpotency in some cases. In fact, it even relax some constraints: most notably, it allows one to consider a broader class of groups. This comes from the fact that the use of the determinant forces to consider only amenable groups, so that the result of the crossed product in Proposition 3.4.3 yields the type II_1 hyperfinite factor. Defining acceptance with a clause on determinant but using nilpotency forced GIRARD to consider languages obtained from finite positive linear combinations of unitaries induced by the group elements. This positivity is needed so that the condition involving the determinant implies nilpotency. Taking directly nilpotency as a marker for acceptance allows us to get rid of this condition, and we could use negative coefficients.

The notion of *normative pair* is a joint restriction on the pair of two subalgebras $(\mathfrak{N}, \mathfrak{D})$ where we pick the representations and the programs. It was defined by GIRARD [59] in order to describe the situations in which an operator in $\mathfrak{D}$ acts uniformly on the set of all representations of a given integer in $\mathfrak{N}$.

Normative Pairs

We will introduce here at the same time the constraint we put on programs so they interact “indifferently” of the representation of the input, and *the* constraint we will use later on. The technical tools needed for the crossed product and in general to grasp the interactions between von Neumann algebras and groups are defined in Section A.2.

Definition 3.4.2 (Normative Pairs). Let $\mathfrak{N}$ and $\mathfrak{D}$ be two subalgebras of a von Neumann algebra $\mathfrak{M}$. The pair $(\mathfrak{N}, \mathfrak{D})$ is a *normative pair* (in $\mathfrak{M}$) if:

- $\mathfrak{N}$ is isomorphic to $\mathfrak{R}$;
- For all $\phi \in M_6(\mathfrak{D})$ and $N_n, N'_n \in M_6(\mathfrak{N})$ two binary representations of n ,

$$\phi N_n \text{ is nilpotent} \Leftrightarrow \phi N'_n \text{ is nilpotent} \quad (3.8)$$

Proposition 3.4.3. Let S be a set and for all $s \in S$, $\mathfrak{N}_s = \mathfrak{R}$. For all group $\mathcal{G}$ and all action α of $\mathcal{G}$ on S , the algebra $\mathfrak{M} = (\bigotimes_{s \in S} \mathfrak{N}_s) \rtimes_{\alpha} \mathcal{G}$ contains a subalgebra generated by $\mathcal{G}$ that we will denote $\mathfrak{G}$. Then for all $s \in S$, the pair $(\mathfrak{N}_s, \mathfrak{G})$ is a normative pair (in $\mathfrak{M}$).

Proof. From the hypotheses, $\mathfrak{N}_s$ is isomorphic to $\mathfrak{R}$. Regarding the second condition (Equation 3.8), we will only show one implication, the other being obtained by symmetry. By Proposition 3.4.2, for

all N_n, N'_n two representation of n in $\mathfrak{N}_s$, there exists a unitary u such that $(\text{Id} \otimes u)N_n(\text{Id} \otimes u)^* = N'_n$. We define $v = \bigotimes_{s \in \mathbb{S}} u$ and π_v the unitary in $\mathfrak{M}$ induced by v . Then π_v commutes with the elements of $\mathfrak{G}$,²⁴ so if there exists $d \in \mathbb{N}$ such that $(\phi N_n)^d = 0$, then $(\phi N'_n)^d = (\phi u N_n u^*)^d = (u \phi N_n u^*)^d = u(\phi N_n)^d u^* = 0$. **Q.E.D.**

Definition 3.4.3 (Observations). Let $(\mathfrak{N}, \mathfrak{G})$ be a normative pair, an *observation* is an operator in $M_6(\mathfrak{G}) \otimes Q$, where Q is a matrix algebra, i.e. $Q = M_s(\mathbb{C})$ for an integer s , called the *algebra of states*.

For the time being, we may just think of Q as a technical trick to have some more space where the action of the group is not needed. Its role will be seen later on. Just remark that from now on, we have to take the tensor product of N_n and Id_Q as an input of an observation.

Definition 3.4.4 (Set “recognized” by an operator). Let $(\mathfrak{N}_0, \mathfrak{G})$ be a normative pair and $\phi \in M_6(\mathfrak{G}) \otimes Q$ be an observation. We define the set of natural numbers:

$$[\phi] = \{n \in \mathbb{N} \mid \phi(N_n \otimes \text{Id}_Q) \text{ is nilpotent, } N_n \in M_6(\mathfrak{N}_0) \text{ any binary representation of } n\}$$

This could be seen as the language recognized by the operator ϕ , but for the time being it is only a mathematical definition.

Definition 3.4.5 (Language decided by a set of observations). Let $(\mathfrak{N}_0, \mathfrak{G})$ be a normative pair and $X \subseteq \bigcup_{i=1}^{\infty} M_6(\mathfrak{G}) \otimes M_i(\mathbb{C})$ be a set of observations. We define the *language decided by X* as the set²⁵:

$$\{X\} = \{[\phi] \mid \phi \in X\}$$

Now we will give the normative pair used as well by GIRARD [59] and by AUBERT and SEILLER [10, 11].

Corollary 3.4.4 (of Proposition 3.4.3). Let $\mathcal{S}$ be the group of finite permutations over $\mathbb{N}$, and for all $n \in \mathbb{N}$, $\mathfrak{N}_n = \mathfrak{R}$. Then $(\mathfrak{N}_0, \mathfrak{G})$ is a normative pair in $(\bigotimes_{n \in \mathbb{N}} \mathfrak{N}_n) \rtimes_{\hat{\alpha}} \mathcal{S}$.

In this particular case, the algebra $(\bigotimes_{n \in \mathbb{N}} \mathfrak{N}_n) \rtimes_{\hat{\alpha}} \mathcal{S}$ is the type II_1 hyperfinite factor. This is one of the reason why GIRARD considered it, as it is then possible to use Fuglede-Kadison determinant. From now on, we will consider this normative pair fixed, and we will study three sets of observations. We first need to define the 1-norm:

Definition 3.4.6 (1-norm). We define the family of norms $\|\cdot\|_1^k$ on $M_k(M_6(\mathfrak{R}))$ by:

$$\|(a_{i,j})_{1 \leq i,j \leq k}\|_1^k = \max_{1 \leq j \leq k} \left(\sum_{i=1}^k \|a_{i,j}\| \right)$$

where $\|a_{i,j}\|$ is the usual norm, i.e. the C^* -algebra norm (Definition A.1.17), on $M_6(\mathfrak{R})$. We will simply write $\|\cdot\|_1$ when the superscript is clear from the context.

²⁴A precise proof of this fact is SEILLER's PhD SEILLER [121, Proposition II.2.5].

²⁵To be consistent with our notation, we should have written them in bold, as $\{\mathbf{X}\}$. But we are going to introduce sets of observations denoted by P_+ , and it could mislead the reader by reminding him/her of $\mathbf{P}$, the polynomial-time hierarchy. Moreover, we wanted to keep the notation adopted in our previous works.

Definition 3.4.7 ($P_{\geq 0}$, P_+ and $P_{+,1}$). An observation $\phi \in M_6(\mathfrak{G}) \otimes M_s(\mathbb{C})$ is said to be *positive* (resp. *boolean*) when for all $0 \leq i, j \leq 6s$, $\phi[i, j]$ is a positive finite linear combination (resp. a finite sum) of unitaries induced by elements of $\mathcal{S}$, i.e. $\phi[i, j] = \sum_{l=0}^{6s} \alpha_l \lambda(g_l)$ with $\alpha_l \geq 0$ (resp. $\alpha_l = 1$).

We then define the following sets of observations:

$$P_{\geq 0} = \{\phi \mid \phi \text{ is a positive observation}\}$$

$$P_+ = \{\phi \mid \phi \text{ is a boolean observation}\}$$

$$P_{+,1} = \{\phi \mid \phi \in P_+ \text{ and } \|\phi\|_1 \leq 1\}$$

We have trivially that $P_{+,1} \subsetneq P_+ \subsetneq P_{\geq 0}$ and $\{P_{+,1}\} \subseteq \{P_+\} \subseteq \{P_{\geq 0}\}$.

3.5 Nilpotency Can Be Decided With Logarithmic Space

We defined a way of representing computation with mathematical objects, but we should now wonder how complex it is to simulate this computation on another device. As usual, the Turing Machine will be the model of computation we chose to prove that the observations are a “reasonable” model of computation.

For the hyperfinite factor is a too complex object to be written on an input tape, we will need to prove that it is possible to work only with the significant part of the operator. This will be the aim of the following subsection, which states and proves a lemma of utmost importance. It is a hard proof that is due to SEILLER and is taken from [10], it uses some mathematical notions that were not defined in this work: we refer to any classical textbook referenced in [Appendix A](#) or to SEILLER [121].

Going back to finite-dimensional

We are going to perform here the reverse operation of the embedding of [Definition 3.4.1](#): we will take the representation of an integer by an infinite operator and enclose it in a finite matrix. And we will do the same for the operator acting on it. Even though the hyperfinite factor $\mathfrak{R}$ is necessary in order to have a uniform representation of integers, it is an algebra of operators acting on an infinite-dimensional Hilbert space. It is therefore not obvious that one can check the nilpotency of such an operator with a Turing Machine, and it is of course not possible in general. However, the following lemma has as a direct consequence that checking the nilpotency of a product ϕN_n where ϕ is an observation in $P_{\geq 0}$ is equivalent to checking the nilpotency of a product of matrices.

Lemma 3.5.1. *We consider the normative pair $(\mathfrak{N}_0, \mathfrak{G})$ defined in [Corollary 3.4.4](#) and denote by $\mathfrak{M}$ the algebra $(\bigotimes_{n \geq 0} \mathfrak{R}) \rtimes \mathcal{S}$. Let N_n be a binary representation of an integer n in $M_6(\mathfrak{N}_0)$ and $\phi \in M_6(\mathfrak{G}) \otimes Q$ be an observation in $P_{\geq 0}$. Then there exist an integer k , an injective morphism $\Psi : M_k(\mathbb{C}) \rightarrow \mathfrak{M}$ and two matrices $M \in M_6(M_k(\mathbb{C}))$ and $\tilde{\phi} \in M_6(M_k(\mathbb{C})) \otimes Q$ such that $\text{Id} \otimes \Psi(M) = (N_n \otimes 1_Q)$ and $\text{Id} \otimes \Psi \otimes \text{Id}_Q(\tilde{\phi}) = \phi$.*

Proof. We denote by n the integer represented by N_n and $R \in M_{6(n+1)}(\mathbb{C})$ its matricial representation. Then there exists a morphism $\theta : M_{n+1}(\mathbb{C}) \rightarrow \mathfrak{R}$ such that $\text{Id} \otimes \theta(R) = N_n$ by [Proposition 3.4.1](#). Composing θ with the inclusion $\mu : M_{n+1}(\mathbb{C}) \rightarrow \bigotimes_{n=0}^N M_{n+1}(\mathbb{C})$, $x \mapsto x \otimes 1 \otimes \cdots \otimes 1$, we get:

$$\text{Id} \otimes \left(\bigotimes_{n=0}^N \theta(\mu(R)) \right) = \overline{N_n} \otimes \underbrace{1 \otimes \cdots \otimes 1}_{N \text{ copies}}$$

where $\overline{N_n}$ is the representation of n in $M_6(\mathbb{C}) \otimes \mathfrak{R}$ (recall the representation N_n in the statement of the lemma is an element of $M_6(\mathbb{C}) \otimes \mathfrak{M}$).

Moreover, since ϕ is an observation, it is contained in the subalgebra induced by the subgroup $\mathcal{S}_N$ where N is a fixed integer,²⁶ i.e. the subalgebra of $\mathcal{S}$ generated by $\{\lambda(\sigma) \mid \sigma \in \mathcal{S}_N\}$. We thus consider the algebra $(\bigotimes_{n=0}^N M_{n+1}(\mathbb{C})) \rtimes \mathcal{S}_N$. It is isomorphic to a matrix algebra $M_k(\mathbb{C})$: the algebra $\bigotimes_{n=0}^N M_{n+1}(\mathbb{C})$ can be represented as an algebra of operators acting on the Hilbert space $\mathbb{C}^{N(n+1)}$, and the crossed product $(\bigotimes_{n=0}^N M_{n+1}(\mathbb{C})) \rtimes \mathcal{S}_N$ is then defined as a subalgebra $\mathfrak{I}$ of the algebra $\mathcal{B}(L^2(\mathcal{S}_N, \mathbb{C}^{(n+1)^N})) \cong M_{(n+1)^N N!}(\mathbb{C})$. We want to show that $(N_n \otimes 1_Q)$ and ϕ are the images of matrices in $\mathfrak{I}$ by an injective morphism Ψ which we still need to define.

Let us denote by α the action of $\mathcal{S}_N$ on $\bigotimes_{n=0}^N M_{n+1}(\mathbb{C})$. We know by [Definition A.2.II](#) that $\mathfrak{I} = (\bigotimes_{n=0}^N M_{n+1}(\mathbb{C})) \rtimes \mathcal{S}_N$ is generated by two families of unitaries:

- $\lambda_\alpha(\sigma)$ where $\sigma \in \mathcal{S}_N$;
- $\pi_\alpha(x)$ where x is an element of $\bigotimes_{n=0}^N M_{n+1}(\mathbb{C})$.

We will denote by γ the action of $\mathcal{S}$ on $\bigotimes_{n=0}^\infty \mathfrak{R}$. Then $\mathfrak{M} = (\bigotimes_{n \geq 0} \mathfrak{R}) \rtimes \mathcal{S}$ is generated by the following families of unitaries:

- $\lambda_\gamma(\sigma)$ for $\sigma \in \mathcal{S}$;
- $\pi_\gamma(x)$ for $x \in \bigotimes_{n \geq 0} \mathfrak{R}$.

As we already recalled, ϕ is an observation in $P_{\geq 0}$ and is thus contained in the subalgebra induced by the subgroup $\mathcal{S}_N$. Moreover, N_n is the image through θ of an element of $M_{n+1}(\mathbb{C})$. Denoting β the action of $\mathcal{S}_N$ on $\bigotimes_{n=0}^N \mathfrak{R}$, the two operators we are interested in are elements of the subalgebra $\mathfrak{I}$ of $\mathfrak{M}$ generated by:

- $\lambda_\beta(\sigma)$ for $\sigma \in \mathcal{S}_N$;
- $\pi_\beta(\bigotimes_{n=0}^N \theta(x))$ for $x \in \bigotimes_{n=0}^N M_{n+1}(\mathbb{C})$.

We recall that ϕ is a matrix whose coefficients are finite positive linear combinations of elements $\lambda_\gamma(\sigma)$ where $\sigma \in \mathcal{S}_N$, i.e. (denoting by k the dimension of the algebra of states):

$$\phi = \left(\sum_{i \in I_{a,b}} \alpha_{a,b}^i \lambda_\gamma(\sigma_{a,b}^i) \right)_{1 \leq a, b \leq 6k}$$

We can therefore associate to ϕ the matrix $\bar{\phi}$ defined as

$$\bar{\phi} = \left(\sum_{i \in I_{a,b}} \alpha_{a,b}^i \lambda_\alpha(\sigma_{a,b}^i) \right)_{1 \leq a, b \leq 6k}$$

We will now use the theorem stating the crossed product algebra does not depend on the chosen representation ([Theorem A.2.I](#)). The algebra $\bigotimes_{n=0}^N \mathfrak{R}$ is represented (faithfully) by the morphism $\pi_\beta \circ \bigotimes_{n=0}^\infty \theta$. We deduce from this that there exists an isomorphism from $\mathfrak{I}$ to the algebra generated by

²⁶That comes from the fact that ϕ “uses” only a finite number of permutations.

the unitaries $\lambda_\beta(\sigma)$ ($\sigma \in \mathcal{S}_N$) and $\pi_\beta \circ \bigotimes_{n=0}^\infty \theta(x)$ ($x \in \bigotimes_{n=0}^N M_{n+1}(\mathbb{C})$). This isomorphism induces an injective morphism ω from $\mathfrak{J}$ into $\mathfrak{J}$ such that:

$$\begin{aligned}\omega(\pi_\alpha(x)) &= \pi_\beta\left(\bigotimes_{n=0}^N \theta(x)\right) \\ \omega(\lambda_\alpha(\sigma)) &= \lambda_\beta(\sigma)\end{aligned}$$

We will denote by ι the inclusion map $\bigotimes_{n=0}^N \mathfrak{R} \subseteq \bigotimes_{n=0}^\infty \mathfrak{R}$ and ν the inclusion map $\mathcal{S}_N \subseteq \mathcal{S}$. We will once again use the same theorem as before, but its application is not as immediate as it was. Let us denote by $\mathcal{S}_N \backslash \mathcal{S}$ the set of the orbits of $\mathcal{S}$ for the action of $\mathcal{S}_N$ by multiplication on the left, and let us choose a representant $\bar{\tau}$ in each of these orbits. Recall the set of orbits is a partition of $\mathcal{S}$ and that $\mathcal{S}_N \times \mathcal{S}_N \backslash \mathcal{S}$ is in bijection with $\mathcal{S}$. As a consequence, the Hilbert space $L^2(\mathcal{S}_N, L^2(\mathcal{S}_N \backslash \mathcal{S}, \bigotimes_{n=0}^\infty \mathbb{H}))$ is unitarily equivalent to $L^2(\mathcal{S}, \bigotimes_{n=0}^\infty \mathbb{H})$. We will therefore represent $\bigotimes_{n=0}^N \mathfrak{R}$ on this Hilbert space and show this representation corresponds to π_γ . For each $x \in \bigotimes_{n=0}^N \mathfrak{R}$, we define $\rho(x)$ by:

$$\rho(x)\xi(\bar{\tau}) = \gamma(\bar{\tau}^{-1})(\iota(x))\xi(\bar{\tau})$$

This representation is obviously faithful. We can then define the crossed product of this representation with the group $\mathcal{S}_N$ on the Hilbert space $L^2(\mathcal{S}_N, L^2(\mathcal{S}_N \backslash \mathcal{S}, \bigotimes_{n=0}^\infty \mathbb{H}))$. The resulting algebra is generated by the operators (in the following, $\xi \in L^2(\mathcal{S}_N, L^2(\mathcal{S}_N \backslash \mathcal{S}, \bigotimes_{n=0}^\infty \mathbb{H}))$):

$$\begin{aligned}\lambda(\nu)\xi(\bar{\tau})(\sigma) &= \xi(\bar{\tau})(\nu^{-1}\sigma) \\ \pi(x)\xi(\bar{\tau})(\sigma) &= \rho(\beta(\sigma^{-1})(x))\xi(\bar{\tau})(\sigma) \\ &= \gamma(\bar{\tau}^{-1})(\gamma(\sigma^{-1})(\iota(x)))\xi(\bar{\tau})(\sigma) \\ &= \gamma(\bar{\tau}^{-1}\sigma^{-1})(\iota(x))\xi(\bar{\tau})(\sigma) \\ &= \gamma((\sigma\bar{\tau})^{-1})(\iota(x))\xi(\bar{\tau})(\sigma)\end{aligned}$$

Through the identification of $L^2(\mathcal{S}_N, L^2(\mathcal{S}_N \backslash \mathcal{S}, \bigotimes_{n=0}^\infty \mathbb{H}))$ and $L^2(\mathcal{S}, \bigotimes_{n=0}^\infty \mathbb{H})$, we therefore get:

$$\begin{aligned}\lambda(\nu)\xi(\sigma\bar{\tau}) &= \xi(\nu^{-1}\sigma\bar{\tau}) \\ &= \lambda_\gamma(\nu)\xi(\sigma\bar{\tau}) \\ \pi(x)\xi(\sigma\bar{\tau}) &= \gamma((\sigma\bar{\tau})^{-1})(\iota(x))\xi(\sigma\bar{\tau}) \\ &= \pi_\gamma(\iota(x))\xi(\sigma\bar{\tau})\end{aligned}$$

Applying [Theorem A.2.1](#) we finally get the existence of an injective morphism ζ from $\mathfrak{J}$ into $\mathfrak{M}$ such that:

$$\begin{aligned}\pi_\beta(x) &\mapsto \pi_\gamma(\iota(x)) \\ \lambda_\beta(\sigma) &\mapsto \lambda_\gamma(\sigma)\end{aligned}$$

[Figure 3.7](#) illustrates the situation. We now define $\Psi : \mathfrak{J} \rightarrow \mathfrak{M}$ by $\Psi = \zeta \circ \omega$. Noticing that $N_n =$

$$\begin{array}{ccccc}
\mathcal{S}_N & \xrightarrow{\lambda_\alpha} & (\bigotimes_{n=0}^N M_{n+1}(\mathbb{C})) \rtimes_\alpha \mathcal{S}_N & \xleftarrow{\pi_\alpha} & \bigotimes_{n=0}^N M_{n+1}(\mathbb{C}) \\
\parallel & & \downarrow \omega & & \downarrow \bigotimes_{n=0}^N \theta \\
\mathcal{S}_N & \xrightarrow{\lambda_\beta} & (\bigotimes_{n=0}^N \mathfrak{A}) \rtimes_\beta \mathcal{S}_N & \xleftarrow{\pi_\beta} & \bigotimes_{n=0}^N \mathfrak{A} \\
\downarrow \subseteq & & \downarrow \zeta & & \downarrow \iota \\
\mathcal{S} & \xrightarrow{\lambda_\gamma} & (\bigotimes_{n \geq 0} \mathfrak{A}) \rtimes_\gamma \mathcal{S} & \xleftarrow{\pi_\gamma} & \bigotimes_{n=0}^\infty \mathfrak{A}
\end{array}$$

Figure 3.7: Representation of the main morphisms defined in the proof of Lemma 3.5.1

$\text{Id}_{M_6(\mathbb{C})} \otimes (\pi_\gamma(\iota \circ \mu(\overline{N_n})))$, we get:

$$\begin{aligned}
\text{Id}_{M_6(\mathbb{C})} \otimes \Psi(M) &= \text{Id}_{M_6(\mathbb{C})} \otimes \Psi(\text{Id} \otimes \pi_\alpha(\text{Id} \otimes \mu)(R)) \\
&= \text{Id}_{M_6(\mathbb{C})} \otimes \pi_\gamma(\iota \circ \bigotimes_{n=0}^N \theta(\mu(R))) \\
&= \text{Id}_{M_6(\mathbb{C})} \otimes \pi_\gamma(\iota(\overline{N_n} \otimes 1 \otimes \cdots \otimes 1)) \\
&= \text{Id}_{M_6(\mathbb{C})} \otimes \pi_\gamma(\iota \circ \mu(\overline{N_n})) \\
&= N_n
\end{aligned}$$

$$\begin{aligned}
\text{Id}_{M_6(\mathbb{C})} \otimes \Psi \otimes \text{Id}_Q(\tilde{\phi}) &= \left(\sum_{i \in I_{a,b}} \alpha_{a,b}^i \Psi(\lambda_\alpha(\sigma_{a,b}^i)) \right)_{1 \leq a,b \leq 6k} \\
&= \left(\sum_{i \in I_{a,b}} \alpha_{a,b}^i \lambda_\gamma(\sigma_{a,b}^i) \right)_{1 \leq a,b \leq 6k} \\
&= \phi
\end{aligned}$$

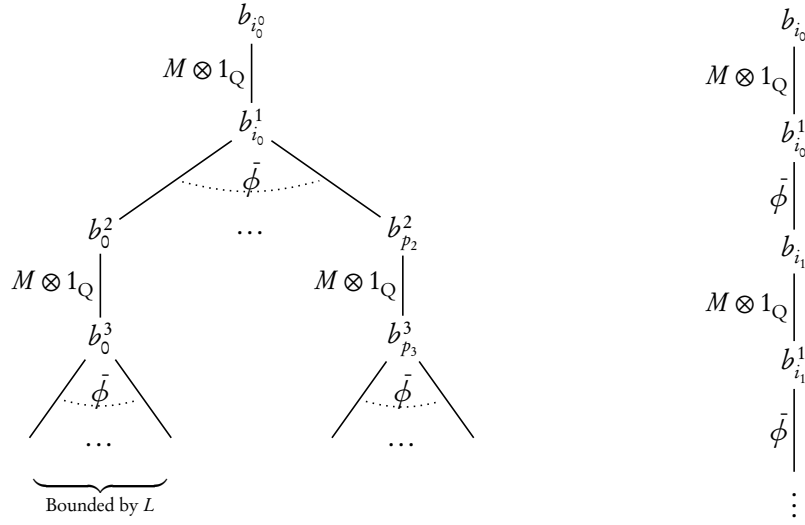
The (injective) morphism Ψ thus satisfies all the required properties.

Q.E.D.

Stated less technically, the previous lemma tells us that no matter the operator ϕ , it always uses only a finite number of elements of $\mathcal{S}$, i.e. of permutations. Therefore, we can identify a finite subgroup $\mathcal{S}_N$ of $\mathcal{S}$, and $m \in \mathbb{N}$ such that $\forall \sigma \in \mathcal{S}, \forall i \geq m, \sigma(i) = i$. It follows that the computation of the iterates of $\phi(N_n \otimes 1_Q)$ takes place in the algebra $M_6(\mathbb{C}) \otimes Q \otimes (\bigotimes_{i=1}^m M_{n+1}(\mathbb{C}) \rtimes \mathcal{S}_N)$, which acts on a *finite* Hilbert space.

Recognizing operators with logarithmic space

This last part is devoted to the study of the complexity of deciding if the product of operators representing a program and an integer is nilpotent. It uses the previous lemma as an evidence that this product can be written as a finite object.



For $\tilde{\phi}$ a matrix obtained from $\phi \in P_{\leq 0}$ and M a representation of an integer, the iteration of $\tilde{\phi}(M \otimes 1_Q)$ may be represented as a tree, where b_i^j denotes the elements of the basis encountered. On the right hand-side, if $\phi \in P_{+,1}$, the computation is simply a line graph.

Figure 3.8: A representation of the computation of an operator

Problem F: Nilpotency of $\phi(N_n \otimes 1_Q)$ for ϕ an operator in $P_{\geq 0}$

Input: $n \in \mathbb{N}$ and $\phi \in P_{\geq 0}$

Output: Accepts iff $\phi(N_n \otimes 1_Q)$ is nilpotent for $N_n \in M_6(\mathfrak{N}_0)$ and $\phi \in M_6(\mathfrak{G}) \otimes Q$.

Theorem 3.5.2. *Problem F is in STA(log, *, $\forall$).*

Proof. The Turing Machine we are going to define cannot take as input ϕ , but we know there is a finite object that behave the same way. Indeed, we know thanks to [Lemma 3.5.1](#) that $\phi(N_n \otimes 1_Q)$ is nilpotent iff $\tilde{\phi}(M \otimes 1_Q)$ is nilpotent. Moreover, the degree of nilpotency of these two products is the same.

Using the bound p given by [Lemma 3.5.1](#), we know that our algebra is

$$M_6(\mathbb{C}) \otimes \underbrace{(M_{n+1}(\mathbb{C}) \otimes \cdots \otimes M_{n+1}(\mathbb{C}))}_{p \text{ copies}} \rtimes \mathcal{S}_N \otimes Q$$

We let K be the dimension of the Hilbert space this algebra is acting on, i.e. for q the dimension of Q , $K = 6(n+1)^p p!q$. The elements of its basis are of the form

$$(\pi, a_1, \dots, a_p; \sigma; e)$$

where $\pi \in \{1, 6\}$ is an element of the basis $(0o, 0i, 1o, 1i, s, e)$ of $M_6(\mathbb{C})$, $a_i \in \{1, \dots, p\}$ are the elements of the basis chosen to represent the integer n , $\sigma \in \mathcal{S}_N$ and e is an element of a basis of Q . When we apply $M \otimes 1_Q$ representing the integer to an element of this basis, we obtain one and only one vector of

the basis $(\pi, a_1, \dots, a_p; \sigma; e)$. The image of a basis element by the observation $\bar{\phi}$ is a finite linear positive combination of $L \in \mathbb{N}$ elements of the basis:

$$\bar{\phi}(\pi, a_1, \dots, a_p; \sigma; e) = \sum_{i=0}^L \alpha_i(\rho, a_{\tau_i(1)}, \dots, a_{\tau_i(p)}; \tau_i \sigma; e_i) \quad (3.9)$$

This combination is finite thanks to [Lemma 3.5.1](#), and it is linear positive by definition of $P_{\geq 0}$ ([Definition 3.4.7](#)).

So we are going to define a universally non-deterministic Turing Machine that will follow the computation on each basis vector thus obtained. The computation can be represented by a tree, as shown in [Figure 3.8](#).

We know that L and the nilpotency degree of $\bar{\phi}(M \otimes 1_Q)$ are both bounded by the dimensions of the underlying space, that is to say by K . Since every coefficient α_i is positive, the matrix is thus nilpotent if and only if every branch of this tree is of length at most K . What ends a branch being of course reaching 0.

This Turing Machine has a logarithmic amount of information to store: the current basis element (i.e. b_{i_0}), the last computed term of the sequence and its index. Every time a branch splits –that is, every time we apply $\bar{\phi}$ –, a non-deterministic transition takes place in order to continue the computation on every branch. Two branches cannot cancel each other, because they have positive coefficients, and the overall product is nilpotent iff every element of its basis goes to 0. As our Turing Machine is “universally non-deterministic”, it can test for nilpotency every element in parallel.

If a branch ends, our Turing Machine halts accepting, if after K applications of $M \otimes 1_Q$ and $\bar{\phi}$ the basis is not null, the Turing Machine starts anew the computation with a “fresh” element of the basis. Once all the elements of the basis has been “tested” and none of them led to nilpotency, the Turing Machine halts rejecting. So the Turing Machine accepts iff there exists an element of the basis such that all branches cancel, that is to say iff $\bar{\phi}(M \otimes 1_Q)$ is nilpotent, and by [Lemma 3.5.1](#) this is equivalent to $\bar{\phi}(N_n \otimes 1_Q)$ nilpotent. Q.E.D.

There is a variant of [Problem F](#), when $\bar{\phi} \in P_{+,1}$, and we define the following problem.

Problem G: Nilpotency of $\bar{\phi}(N_n \otimes 1_Q)$ for $\bar{\phi}$ an operator in $P_{+,1}$

Input: $n \in \mathbb{N}$ and $\bar{\phi} \in P_{+,1}$

Output: Accepts iff $\bar{\phi}(N_n \otimes 1_Q)$ is nilpotent for $N_n \in M_6(\mathfrak{N}_0)$ and $\bar{\phi} \in M_6(\mathfrak{G}) \otimes Q$.

Theorem 3.5.3. *Problem G is in STA(log, *, d).*

Proof. In the case when $\bar{\phi}$ is obtained from an observation $\phi \in P_{+,1}$. In this particular case, [Equation 3.9](#) becomes:

$$\bar{\phi}(\pi, a_1, \dots, a_p; \sigma; e) = (\rho, a_{\tau_{b_0}(1)}, \dots, a_{\tau_{b_0}(p)}; \tau_{b_0} \sigma; e_{b_0}) \quad (3.10)$$

Indeed, writing $\bar{\phi}$ as a $\text{Card}(Q) \times \text{Card}(Q)$ matrix (here Q denotes a basis of Q , which is a matrix algebra from [Definition 3.4.3](#) of observations), we use that $\|\bar{\phi}\|_1 \leq 1$ to deduce that for all $q \in Q$, $\sum_{q' \in Q} \|\bar{\phi}[q, q']\| \leq 1$.

From the equation:

$$\bar{\phi}(\pi, a_1, \dots, a_p; \sigma; e) = \sum_{i=0}^l \alpha_i(\rho, a_{\tau_i(1)}, \dots, a_{\tau_i(p)}; \tau_i \sigma; e_i)$$

we have $\|\bar{\phi}[e, e_i]\| = 1$ ($i = 1, \dots, l$), and thus $\sum_{e' \in Q} \|\bar{\phi}[e, e']\| \geq l$. Therefore, $l \leq 1$, and we refer to Figure 3.8 for a graphical representation.

Notice that the nilpotency degree of $\bar{\phi}M$ is again bounded by K (the dimension of the Hilbert space the algebra is acting on). One can thus easily define a deterministic Turing Machine that takes the basis elements one after the other and compute the sequence $b_{i_0}, b_{i_1}, \dots$: if the sequence stops before the K th step, the machine starts with the next basis element as b_{i_0} , and if the sequence did not stop after K step it means the matrix was not nilpotent. **Q.E.D.**

Corollary 3.5.4.

$$\{P_{\geq 0}\} \subseteq \text{co-NL}$$

$$\{P_{+,1}\} \subseteq \mathbf{L}$$

Proof. From Lemma 3.5.1, given an observation and an integer, we know there is a finite product of matrix that is nilpotent iff the product of this observation and integer is nilpotent. Then, by Theorem 3.5.2 and Theorem 3.5.3, we know that a log-space Alternating Turing Machine can decide if this product is nilpotent, i.e. if the observation is going to accept or reject the input. If this observation is of 1-norm strictly greater than 1, the Alternating Turing Machine has to be universally non-deterministic, to test in parallel every element of the basis. If this observation is of 1-norm less or equal to 1, the Alternating Turing Machine only has to test the elements of the basis one after the other, therefore it can be deterministic. **Q.E.D.**

Results and perspectives

What have we done?

We showed how to represent data and programs as mathematical objects. We first used the Church encoding and the Curry-Howard isomorphism to represent data—in this case integers—as proofs. Then, by using the proof net representation, we showed how it was possible to represent those proofs as matrices. As this representation was not uniform, i.e. the size of the matrix depends on the size of the integer, we showed how the matrices could be embedded in the Π_1 hyperfinite factor. The motivations to use this mathematical object was exposed in the first section.

There is infinitely many ways of representing integers as operators in a von Neumann algebra, but we managed to define a framework—normative pairs—where integers and programs interact nicely, i.e. where the programs are indifferent to the particular representation chosen. We used the most convenient normative pair, i.e. the one that internalise permutations over $\mathbb{N}$. We chose to take nilpotency as a criteria of acceptance rather than the cancellation of the determinant, and that allows us to consider broader classes of groups.

At last by using a lemma proven by Thomas SEILLER, we showed that it was possible to come back to finite objects and to decide their nilpotency with logarithmic space resources. This algebra has for the time no precise computational meaning, but we introduced a norm that somehow represent deterministic computation, the “basic” computation being widely non-deterministic, and to be more precise universally non-deterministic.

What could be done?

The purely mathematical way of modelling computation we defined in this chapter is a brand new and exciting field of research. Our first task will be to understand its computational nature, and that will

be one of the purposes of the next chapter, by proving that a specific kind of “pointer machine” can be reasonably simulated by operators.

And yet, there are already numerous hints toward extensions or modification of this framework:

1. We could try to embed other kinds of data in the von Neumann algebras, to have several constructors. As long as a data can be represented by a proof, we can follow exactly the same construction and obtain a representation in the hyperfinite factor. For instance, any finite list of objects of type U can be typed with $\forall X X \rightarrow ((U \rightarrow (X \rightarrow X)) \rightarrow X)$ in system F [62, pp. 90 sqq.]. This could take place in a general movement to go back to logical considerations.
2. As we sketched in Section 3.4, we could try to take the crossed-product of a von Neumann algebra with other groups. A good candidate could be “l’algèbre des boîtes”, introduced by GIRARD [56, pp. 506 sqq.]. That would be really nice if we could prove some property like “if $\mathcal{G}$ is a strict subgroup of $\mathcal{H}$, then the set of elements recognized by a crossed product with $\mathcal{G}$ is strictly included in the set of elements recognized by a crossed product with $\mathcal{H}$.” That could lead to separation results, but that the construction maintains strict inclusion could be hard to prove.
3. We could also play on another level with our normative pair, by allowing negative coefficients. In this setting, there could be cancellation between different “branches” of a computation, thus giving to our decomposition of the basis the structure of a graph —instead of the tree we saw in Figure 3.8. Such sharing in the computation reminds of classes of complexity stronger than log-space computation.
4. Another perspective could be to try to get rid of von Neumann algebras.²⁷ We could define observations that act only on a fixed size of input and stay on a matricial level. We should be really cautious about duplication in such a setting, but as well as one could give a bound on time for a space-bounded Turing Machine, maybe we could guarantee that an operator would not need matrices bigger than a pre-computed bound. That would probably need a uniformity similar to the one defined for the Boolean circuits (recall Definition 2.1.4), i.e. to define how to build a family of observations.
5. A last perspective would be to consider some advanced works in von Neumann algebra, which is a very active field of research. For instance, FULMAN [43] developed a construction similar to the crossed product of a von Neumann algebra by a group, the crossed product of an abelian von Neumann algebra by an equivalence relation. Surely a discussion with an expert on those topics could leads to numerous other ideas!

²⁷Since the writing of this memoir, GIRARD developed a new approach on Geometry of Interaction, using unification techniques [60]. A large part of the work presented in Chapter 3 and Chapter 4 has been rephrased in terms of “unification algebra” in a joint work with Marc BAGNOL [9].

Nuances in Pointers Machineries

Contents

3.1	Complexity and Geometry of Interaction	58
	Past approaches	
	Complexity with operator algebra	
3.2	First Taste: Tallies as Matrices	60
	First representation of the unary integers, without the contraction rule	
	Unary integers with the contraction rule	
3.3	Binary Integers	68
	Binaries as proofs	
	Binaries as matrices	
	Some remarkable properties of the matricial representation	
3.4	How to Compute	76
	Binaries as operators	
	Where programs and integers interact	
	Normative Pairs	
3.5	Nilpotency Can Be Decided With Logarithmic Space	83
	Going back to finite-dimensional	
	Recognizing operators with logarithmic space	

OUR goal in this chapter will be to build a computational model that accounts for the computation of the operators (or observations) defined in the previous chapter. For reasons that should be clear after reading [Section 4.4](#), the computation of operators may be seen as a wander in the structure of the binary word that cannot modify it, but that may “save” values in registers. It is quite natural to think of *pointers* in this setting, and to GIRARD [59] it went without saying that the stroll of some “fingers” was the log-space-computation.

However, we saw that speaking of “the” log-space computation could be misleading: recall ([Definition 1.1.11](#)) that depending on the bounds on time and alternation, an Alternating Turing Machine with

only a logarithmic amount of space can characterise from NC^1 up to P ! So another approach—more reasonable—is to start from this idea of fingers wandering in an input and to wonder to what model and what complexity class it could correspond. That makes us dive into another vivid field of research, closer to computer science, that deals with pointers, and we will see that depending on the kind of pointers we manipulate and the way we manipulate them, we can capture more, less, or exactly L .

We propose here a chapter mid-way between a brief summary of the “pointer machines” that were defined and an illustration of the flexibility of a model of computation—Finite Automaton—, but both will be necessary to achieve our goal, i.e. to define *Non-Deterministic Pointer Machines* as the model of computation that corresponds to operators. We will prove that claim thanks to a low-level simulation: the proof is mainly based on the definition of the encoding of instructions (read a value, move a pointer, change the state) as operators. This section will be quite long, for we have to recall the framework of Chapter 3 and introduce notations to ease the simulation. We moreover handle the case of deterministic computation.

The first two sections were greatly inspired by some discussions with Arnaud DURAND, and Paulin JACOBÉ DE NAUROIS gave me some precious hints during the redaction of this chapter. The first part of the following subsection witness a work led with Margherita ZORZI and Stefano GUERRINI that unfortunately turned out to be a deadlock. Some elements of this chapter, especially the last two sections, is a joint work [10, 11] with Thomas SEILLER and constitute our main contribution. Some elements, as Theorem 4.2.2, are presented in an innovative way.

It could be helpful to keep in mind some notions on descriptive complexity to fully grasp the first section of this chapter,¹ but apart from that, everything was defined in Chapter 1 and in Chapter 3. Some parts of the second section may be really redundant to a complexity expert, but we felt it could be advantageous to recall some fundamentals in this topic.

4.1 A First Round of Pointers Machines

We will not make the inventory of the pointer machines that were defined, for they are *really* numerous.² BEN-AMRAM [16] made it clear that the heading of “pointer machines” was really misleading, for it covers at the same time “pointer algorithms” and “abstract machine models”, both atomistic and high-level. We will apply his advice and call the atomistic pointer machines model we are going to work with ... pointers machines! We will nevertheless blur this distinction, for the subjects we are going to evoke are sometimes midway between models of computation, programming languages and descriptive complexity: it comes from the “fact that it is natural to implement a Pointer Algorithm using the capabilities of [SMM or LISP machines].” ([16, p. 94])

Instead, this section will go against the intuitive idea that “pointers are log-space”, by showing that depending on the precise definition of the machine that manipulates “pointers” (whose definition also fluctuate), we can characterise less (with JAG and PURPLE) or more (with SMM and KUM) than L . The exact characterization of L by a pointer machine will be postponed to the next section.

This journey towards the understanding of the “pointer complexity” will help us to define a subtle variation of Finite Automata that will end up characterizing the kind of computation we mimic with the operators of Chapter 3.

¹And the reader may with benefit take a look at two classical textbooks, IMMERMANN [79] and EBBINGHAUS and FLUM [40].

²We could only mention TARJAN [129]’s reference machine as one of the first model refereed to as “pointer machine”.

JAGs and PURPLE: on the limitations of pebbles and pure pointers

We begin with two models, strongly linked, that cannot decide a problem that was proven to belong to L .

Jumping Automaton for Graphs (JAG) were introduced by COOK and RACKOFF [29] to try to demonstrate a lower-bound for *maze-threadability* problem, a variation on *STConn* (recall Problem B). This model takes a graph as input, and moves pebbles, or tokens, along its edges. It has a finite state control, the ability to know if two pebbles are at the same vertex and to make one pebble “jump” to another. This model can easily be considered as a strict restriction of Turing Machine, it illustrates once again the importance of graphs for the log-space computation. It manipulates pointers in the sense that those pebbles are only location, addresses, in the input graph, they cannot be used to store any other kind of information. This model was designed to solve accessibility problems, but numerous modifications were proposed to handle other problems.

And yet, HOFMANN and SCHÖPP [73] thought that this approach could be led further, for “neither JAGs nor [Deterministic Transitive Closure (DTC)]-logic adequately formalise the intuitive concept of “using a constant number of graph variables only”. This work is also greatly inspired by CAI, FÜRER, and IMMERMANN [22], who proved that First-Order Logic (FO) with least fixed point operator and counting, but without ordering, failed to capture P . So they define PURPLE, in order to make formal the idea that “log-space is more than ‘a constant number of pointers’”.

PURPLE is a *while* programming language, without quantifiers loop, and it is “pure” in the sense that pointers are abstract values, without internal structure. PURPLE’s input are *partially ordered* graphs, for a precise reason: if there was a total order on graphs, then one could encode information in it. Without going further into details, we can mention that on ordered structures, FO and DTC corresponds to L , but it is not true on partially ordered structures. They tried to prove lower bounds and this approach was thrilling, for they proved that

- PURPLE cannot solve “the number of nodes of a graph is a power of 2” nor *USTConn* [73] which are in L . This remains true even if we add to PURPLE iteration (the quantifier loop) (this was a new lower bound, the purpose of [72]).
- With counting, PURPLE can decide *USTConn*.
- By introducing some elements of non-determinism, HOFMANN, RAMYAA, and SCHÖPP [71] tries to prove that $L \subsetneq PTIME$. By doing so, they managed to show that even with counting, a non-deterministic version of PURPLE could not solve a L -complete problem, tree isomorphism.

Those numerous variations on PURPLE help to provide new lower bounds, and it is inspired by techniques of descriptive complexity. In this perspective, one can easily add or remove a single “ingredient to our recipe”, the logic under study, and try to establish new bounds. The interested reader may for instance refer to the work of GRÄDEL and MCCOLM [66], where it is proven that transitive closure logic is strictly more powerful than DTC.

We tried with Margherita ZORZI and Stefano GUERRINI to link together PURPLE and GIRARD’s approach. Unfortunately, we came to the conclusion that our functional setting could not match the imperative setting of PURPLE, for we could encode structured information in our finite control.

KUM and SMM: questioning the meaning of *log-space*

The Kolmogorov Uspenskiĭ Machine (KUM) was first presented in Russian [89], but the English translation[90] of this article is better known. This model of computation is a “pointer machine” in the sense that instead of operating on registers —displayed as a tape— it manipulates a graph, not only by reading it —as for PURPLE—, but by making it evolves.

The Storage Modification Machine (SMM), introduced by SCHÖNHAGE [116], was defined independently, but can be conceived as a special kind of KUM where the graph manipulated is directed, whereas it is undirected for the KUM. We give below an informal definition of the SMM.

Definition 4.1.1 (SMM). A SMM is composed from a finite control given as a program and a dynamic Δ -structure for Δ an alphabet. The machine reads an input string and writes on an output string, and uses as a storage a dynamic labelled graph (X, a, p) , where

- X is a finite set of nodes,
- $a \in X$ is the center (the active node),
- $p = (p_\alpha \mid \alpha \in \Delta)$ is the family of pointer mappings $p_\alpha : X \rightarrow X$

To write $p_\alpha(x) = y$ means that the pointer with label α originating from x goes to y . We define $p^* : \Delta^* \rightarrow X$ recursively:

$$\begin{aligned} p^*(\epsilon) &= a, \text{ where } \epsilon \text{ is the empty word} \\ p^*(W\alpha) &= p_\alpha(p^*(W)) \text{ for all } \alpha \in \Delta, W \in \Delta^* \end{aligned}$$

A SMM manipulates a graph by moving its “center” in it, by merging and creating nodes, by adding or removing pointers, i.e. labelled edges. It manipulates pointers in the sense that p expresses *addresses*, it encodes paths from the center to any node.

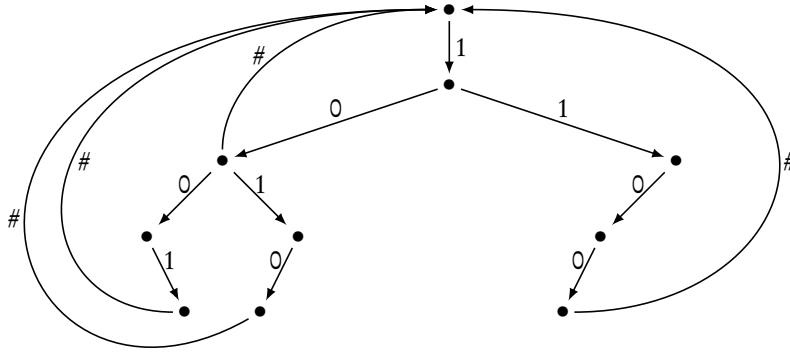
One simple example should show how to compute with this kind of machine.

Example 4.1.1 (Sort with SMM). We want to build a simple program that sort integers. Suppose we want to sort 9, 12, 10 and 2, the input will be #1001#1100#1010#10#. The SMM will proceed as follows: every time the bit currently read is a 0 (resp. a 1), it tests for an edge labelled with 0 (resp. 1) leaving its actual center. If it exists, it follows it, elsewhere it creates a node, add an edge labelled with 0 (resp. 1) and follows it. When it reads a #, it goes back to its initial center. For the input we gave, when the SMM reaches the end of the input tape, the graph looks as in Figure 4.1

Then, the results is read simply by a depth first search, following in priority the edges labelled by #, 0, and then 1. This gives us for the current example : 10, 1001, 1010, 1100, i.e. the integers sorted.

The algorithm we gave is *linear* in the size of the input if we take the “uniform time measure”, *a.k.a.* one transition count for one unit of time. It might look *too efficient*, for the best comparison-based sorting algorithms are $\mathcal{O}(n \log(n))$, and $\mathcal{O}(\log^2(n))$ in parallel.³ All the same, it was proven [116, Theorem 6.1, p. 503] that a SMM could perform integer-multiplication in linear time, whereas the best known algorithm for other models needs $\mathcal{O}(n \log(n))$ steps.

³Of course, it is a different story if we consider *Counting Sort* or *Bucket Sort* algorithms.



Remember we are sorting #1001#1100#1010#10#.

Figure 4.1: How a KUM sorts integers

It is the same regarding space: it seems reasonable to consider that the number of nodes is the measure of space-complexity of a SMM. But on a Δ -structure with n nodes and if $\text{Card}(\Delta) = k$, one can encode $\mathcal{O}(nk \log(n))$ bits of information, whereas a Turing Machine with n cells on alphabet Δ can encode $\mathcal{O}(n \log(k))$ bits. This results follows from a combinatorial analysis [116, p. 492] of the number of graphs of n nodes with edges labelled by elements of Δ .

In this setting, “one should be aware of the fact that for such a model **log-space** has become a different concept.” (VAN EMDE BOAS [131, p. 106]) SMM and KUM does not belong to the class of “reasonable models of computation”, for their time- and space-complexity —or at least the measure we defined on them— do not seem realistic.⁴ They are a good way to express an algorithm, a sort of generalized or universal pointer machine that was rather used as a way of conceiving the problems and the algorithms than a model of computation actively used.

Yet, the ICC framework gave some new perspectives to that model with the recent work of LEIVANT and MARION [97]. They develop a program with a dynamic data-structure, that support the creation, deletion and update of vertices and edges. Thanks to ramification, they managed to enforce a polynomial time-complexity on that structure.

It seems really complicated to see the computational behaviour of observations as a rewriting of a graph. We were addressing the notion of pointer rather because our framework does not seem to have the ability to write, but only to parse. So no matter how interesting that model is, we will move to another, more suited. The reader who would like to learn more about this subject should look at the special issue of the Journal for Universal Computer Science [21], which is devoted to the celebration of the ten years of GUREVICH’s Abstract State Machines.

⁴This is a matter of taste: one could perfectly say that those models are the “reasonable” ones, and that the other models are too bureaucratic, inefficient, to handle smoothly complex algorithms based on pointers.

4.2 Finite Automata and L

We will recall the definition of a classical model of computation, Multi-Head Finite Automaton (FA), a read-only variant of the multi-head Turing Machines.⁵ This model was introduced in the early seventies by HARTMANIS [69] and provided many interesting results. Their nice characterization of **L** and **NL** (Theorem 4.2.4) and the similarities they share with Non-Deterministic Pointer Machines motivate their presence here in this small tutorial: even if they are called “(read-only) heads”, they are exactly what we mean when we speak of “pointers”. We will play a bit with this model of computation, take some times to understand how to compute with pointers and illustrate the flexibility of this model.

All the definitions and properties that follow are classical results, stated in a uniform way whereas the references sometimes mixes the notations. They are mostly borrowed to the recently published HOLZER, KUTRIB, and MALCHER [75], which is a nice overview of the main results and open problems of the theory. Whereas the number of heads is an important criteria from a complexity point of view (in fact, it is the only relevant one), and even if separations results are known regarding that criteria (cf. Theorem 4.2.1), we won’t count precisely the number of heads (the reason should be clear after looking at Theorem 4.2.4) in the modifications that follows.

Presentation and variations

Definition 4.2.1 (Non-Deterministic Multi-Head Finite Automaton (NFA)). For $k \in \mathbb{N}^*$, a *non-deterministic two-ways k -heads finite automaton* (NFA(k)) is a tuple $M = \{S, A, k, \triangleright, \triangleleft, s_0, F, \sigma\}$ where:

- S is the finite set of *states*;
- A is the *alphabet*;
- k is the number of heads;
- $\triangleright$ and $\triangleleft$ are the *left* and *right endmarkers*, $\triangleright, \triangleleft \notin A$;
- $s_0 \in S$ is the *initial state*;
- $F \subseteq S$ is the set of *accepting states*;
- $\sigma \subseteq (S \times (A \cup \{\triangleright, \triangleleft\})^k) \times (S \times \{-1, 0, +1\}^k)$ is the *transition relation*, where -1 means to move the head one square to the left, 0 means to keep the head on the current square and 1 means to move it one square to the right.

Moreover, for all $a_1, \dots, a_k \in A^k$, $1 \leq i \leq k$ and $m_i \in \{-1, 0, +1\}$, whenever⁶

$$(s, (a_1, \dots, a_k)) \sigma (s', (m_1, \dots, m_k))$$

then $a_i = \triangleright$ implies $m_i \in \{0, +1\}$, and $a_i = \triangleleft$ implies $m_i \in \{-1, 0\}$. That is to say that the heads cannot move beyond the endmarkers.

⁵Although this model was first defined as a variant of Finite Automaton, we thought it would be more convenient to introduce it as a variation of Turing Machine. We will nonetheless use some of their properties as a Finite Automaton, and will refer to the literature of this field. For an excellent overview of this theory, we refer for any classical textbook, like the one by HOPCROFT, MOTWANI, and ULLMAN [76].

⁶Of course, one may see σ as a partial function from $S \times (A \cup \{\triangleright, \triangleleft\})^k$ to $\mathcal{P}_{\text{fin}}(S \times \{-1, 0, +1\}^k)$, and that justifies this notation that we will use from now on.

One defines configurations and transitions in a classical way. Let $M \in \text{NFA}(k)$ and $w \in A^n$ an input of length n . We say that M *accepts* w if M starts in state s_0 , with $\triangleright w \triangleleft$ written on its input tape and all of its heads on $\triangleright$, and if after a finite number of transitions at least one branch of M reaches a state belonging to F and cannot perform any transition. We say that M *always halts* if for all input all branches of M reach after a finite number of transitions a configuration such that no transition may be applied. Remark that if no transition may be applied and the current state is not in F , this branch rejects the input. If all branches reject the input, that is, they all halt in a state not in F , M rejects the input.

If σ is functional, M is said to be *deterministic* and we write $\text{DFA}(k)$ the set of deterministic two-way k heads automata (DFA). We will write Finite Automata (FA) in all generality, if we don't want to be specific about the (non-)determinacy of our Finite Automata. There are numerous other variation of the FA that we did not mention. Among them, the 1-way (Non-Deterministic) Finite Automata, which has exactly the same definition as the NFA, except that the heads can only go from left to right, i.e. -1 is not in the set of movements. This model is however interesting, for it provided some separation results and helped to be even more precise in the study of the computational power of the FA (we refer to the following subsection for that topic).

We will denote as usual by $\mathcal{L}(M)$ the language accepted by M . We will denote by **DFA** and **NFA** the sets of languages decided by respectively the sets $\cup_{k \in \mathbb{N}^+} \text{DFA}(k)$ and $\cup_{k \in \mathbb{N}^+} \text{NFA}(k)$. The weight of the font should help the reader to get the context: a machine belongs to $\text{NFA}(k)$ and a problem to **NFA**.

Now we will assert several claims and give ideas of proofs. They are quite classical and some of their equivalent for Turing Machines were asserted without any proof, but we felt it should be made precise in this case.

Remember for instance that due to its definition, it is not trivial at all that for all $M \in \text{DFA}(k)$ there exists $M' \in \text{DFA}(k)$ such that M accepts iff M' rejects! It was an open problem for HARTMANIS [69], denoting by C_D^k “the set of all languages accepted by deterministic k -head automata”: “It would be also interesting to determine whether for every $A \in C_D^k$ also $\bar{A} \in C_D^k$. Again we do not know the answer to this problem. If C_D^k should not be closed under complementation then it would be interesting to find out how many additional heads are needed to obtain the complements of all sets in C_D^k . We know that $2k$ heads suffice but we suspect that a considerably smaller number of heads is really required to accept the complements.” We won't bother answering this question here: all we need to know is that **DFA = co-DFA**.

Before we develop some basic computation with our FA, we define some syntactical sugar:

- For every alphabet A , we define A' to be $A \cup \{\triangleright, \triangleleft\}$.
- We write $*$ for any symbol in A' , $[\text{not } a]$ for any symbol in $A' \setminus \{a\}$.
- We will sometimes do as if the heads could make several movements in one transition, i.e. take $\sigma \subseteq (S \times (A'^k) \times (S \times \{q\}^k))$ with $q \in \mathbb{Z}$. It is only a shorthand, for it is really easy to modify a FA that uses this feature to a one that does not.
- For all head H , we write $\#H$ the distance (i.e. the number of cells) between $\triangleright$ and H .

This last point will be used to show how some information can be encoded in *the position* of the heads: what will matter won't be any more the value pointed, but its position in the input string.

Every of our claims will apply as well to the deterministic and the non-deterministic cases, and the resources needed to process the translation and the overhead will be studied only for the first case, the proof being always the same.

CLAIM 1: *We can always take the alphabet A to be $\{0, 1\}$.*

Proof. Formally, we prove that for all $k \in \mathbb{N}^*$, all $M \in \text{FA}(k)$ with alphabet $A \neq \{0, 1\}$ there exists $M' \in \text{FA}(k)$ with $B = \{0, 1\}$ and $f : A^* \rightarrow B^*$ such that $f(\mathcal{L}(M)) = \mathcal{L}(M')$. We will prove moreover that M' simulates M in real-time.

Let p be such that $2^{p-1} < \text{Card}(A) \leq 2^p$. If $p = 1$, then f is simply a renaming of A . If $p > 1$, we encode any symbol of A with p bits in B . This is simply done by associating to every symbol of A a binary integer expressed with p bits.

Then, for all $a_1, \dots, a_k \in A^k$ and $(m^1, \dots, m^k) \in \{-1, 0, +1\}^k$, every transition of M is of the shape

$$(s, (a^1, \dots, a^k)) \sigma (s', (m^1, \dots, m^k)) \quad (4.1)$$

For all such transitions, we introduce p states $s_{a_1^i \dots a_1^k}, \dots, s_{a_1^i a_2^i \dots a_p^i a_2^k \dots a_p^k}$ where a_w^j denotes the w th bit of the binary encoding of a^j , i.e. $a_i \equiv a_1^i \dots a_p^i$.

We then translate the transition of Equation 4.1 by

$$\begin{aligned} & (s, (a_1^i, \dots, a_1^k)) \sigma' (s_{a_1^i \dots a_1^k}, (+1, \dots, +1)) \\ & (s_{a_1^i \dots a_1^k}, (a_2^i, \dots, a_2^k)) \sigma' (s_{a_1^i a_2^i \dots a_1^k a_2^k}, (+1, \dots, +1)) \\ & \vdots \\ & (s_{a_1^i a_2^i \dots a_p^i a_2^k \dots a_p^k}, (a_m^i, \dots, a_m^k)) \sigma' (s', (m^1, \dots, m^k)) \end{aligned}$$

$$\text{where } m'^j = \begin{cases} -(2 \times p) & \text{if } m^j = -1 \\ -p & \text{if } m^j = 0 \\ +1 & \text{if } m^j = +1 \end{cases}$$

We can easily check that M accepts an entry $w \in A^n$ iff M' accepts $f(w) \in \{0, 1\}^{n \times \log(p)}$. The number of states of M' grows quite violently compared to M , but M' does not need more heads and simulates M in real time, i.e. with a linear overhead. **Q.E.D.**

It is worth mentioning that two letters- and three letters-alphabets are not always equivalent, as proven by KOZEN [91], for instance when we study their properties as context-free languages. We won't study such properties here, but we have to pay attention to the complexity of f and the resources needed to describe M' given n and the description of M . The encoding f is really classic, and can be performed by a constant-depth circuit. The description of M' is mainly obtained by inserting routines in the description of M . With a log-time bounded ATM, one can guess in parallel the description of M , label every state with the k^p different values that may be read, and modify the transition function accordingly. This translation can then be performed with NC^1 resources, using the power of alternation of the ATM. The following claims admit the same bounds on the complexity of the translation.

CLAIM 2: *It is possible to move only one head at every transition.*

Proof. Formally, we prove that for all $k \in \mathbb{N}^*$, for all $M \in \text{FA}(k)$, there exists $M' \in \text{FA}(k)$ such that $\mathcal{L}(M) = \mathcal{L}(M')$ and such that for all transition in M' ,

$$(s_1, (a_1, \dots, a_k)) \sigma' (s_2, (m_1, \dots, m_k))$$

for all $1 \leq i \leq k$, only one among m_i is different from 0.

Every translation in M

$$(s_1, (a_1, \dots, a_k)) \sigma (s_2, (m_1, \dots, m_k))$$

can be translated, for s_1^i k “fresh” states, i.e. not in S , by

$$\begin{aligned} & (s_1, (a_1, \dots, a_k)) \sigma' (s_1^1, (m_1, 0, \dots, 0)) \\ & (s_1^j, (*, \dots, *)) \sigma' (s_1^{j+1}, (0, \dots, 0, m_j, 0, \dots, 0)) \quad (\text{For } 1 \leq j < k) \\ & (s_1^k, (*, \dots, *)) \sigma' (s_2, (0, \dots, 0, m_k)) \end{aligned}$$

Q.E.D.

Definition 4.2.2 (Sensing heads). We say that a FA has *sensing heads* if it has states $s_{i=j?}$ such that

$$(s_{i=j?}, (a_1, \dots, a_k)) \sigma \begin{cases} (s_1, (m_1, \dots, m_k)) & \text{if } \#H_i = \#H_j \\ (s_2, (m'_1, \dots, m'_k)) & \text{elsewhere} \end{cases} \quad (4.2)$$

In other words, our FA can “sense” if two heads are at the same cell and act accordingly. Note that it has nothing to do with non-determinacy: only one among the two transition of Equation 4.2 is applied.

CLAIM 3: *One additional head is sufficient to mimic the sensing heads ability.*

Proof. We prove that for all $M \in \text{FA}(k)$ with sensing heads, there exists $M' \in \text{FA}(k+1)$ without sensing heads such that $\mathcal{L}(M) = \mathcal{L}(M')$.

The simulation needs an additional head H_{k+1} that is on the beginning of the tape. Every translation of the form of Equation 4.2 is translated thanks to two fresh states $s_{i=j}$ and $s_{i \neq j}$, with the following routine⁷:

$$\begin{aligned} & (s_{i=j?}, (\dots, a_i, \dots, a_j, \dots, \triangleright)) \sigma' (s_{i=j?}, (\dots, -1, \dots, -1, \dots, +1)) \\ & (s_{i=j?}, (\dots, [\text{not } \triangleright], \dots, [\text{not } \triangleright], \dots, *)) \sigma' (s_{i=j?}, (\dots, -1, \dots, -1, \dots, +1)) \\ & \left. \begin{aligned} & (s_{i=j?}, (\dots, \triangleright, \dots, [\text{not } \triangleright], \dots, *)) \\ & (s_{i=j?}, (\dots, [\text{not } \triangleright], \dots, \triangleright, \dots, *)) \end{aligned} \right\} \sigma' (s_{i \neq j}, (\dots, +1, \dots, +1, \dots, -1)) \\ & (s_{i=j?}, (\dots, \triangleright, \dots, \triangleright, \dots, *)) \sigma' (s_{i=j}, (\dots, +1, \dots, +1, \dots, -1)) \\ & (s_{i=j}, (\dots, *, \dots, *, \dots, [\text{not } \triangleright])) \sigma' (s_{i=j}, (\dots, +1, \dots, +1, \dots, -1)) \\ & (s_{i \neq j}, (\dots, *, \dots, *, \dots, [\text{not } \triangleright])) \sigma' (s_{i \neq j}, (\dots, +1, \dots, +1, \dots, -1)) \\ & (s_{i=j}, (\dots, *, \dots, *, \dots, \triangleright)) \sigma' (s_1, (m_1, \dots, m_k, 0)) \\ & (s_{i \neq j}, (\dots, *, \dots, h_j, \dots, \triangleright)) \sigma' (s_2, (m'_1, \dots, m'_k, 0)) \end{aligned}$$

⁷Where the $\dots$ are to be read as $*$ on the left hand side, and as 0 on the right hand side of σ' .

The symbol $\triangleright$ has the role of a buttress, the configurations store if we had $\#H_i = \#H_j$ and we can check that at the end of the test H_i , H_j and H_{k+1} are back to their position. We also need to modify all the other transition to add “whatever H_{k+1} reads” and “ H_{k+1} does not move”.

Every time M performs a comparison, M' have to perform at most $2 \times n$ transitions. **Q.E.D.**

We considered in Section 1.1 that it was trivial to define a Turing Machine that always halts from a Turing Machine that may loop, provided a bound, and we will prove the same for FA. We felt it necessary to give a proof, for it is amusing (the heads acts like hands of a clock) and not evident that without a writing ability such mechanism could be developed.

CLAIM 4: *We can consider only FA that always stops.*

Proof. We will prove that for all $M \in \text{FA}(k)$, there exists $M' \in \text{FA}(k')$ such that M' always halt and $\mathcal{L}(M) = \mathcal{L}(M')$.

Given an input $w \in A^n$, we know that the number of configurations of M is bounded by $\text{Card}(S) \times (\text{Card}(A'))^k \times (n+2)^k$, lets $d \in \mathbb{N}$ be such that the number of configurations is inferior to n^d . We will construct M' so that it halts rejecting after performing more than this number of transitions.

We set $k' = k + d$, and we construct M' so that its k first heads act as the heads of M , and the d heads act as the hands of a clock, going back and forth between the two endmarkers. We set

$$S' = \{s \times \{\rightarrow, \leftarrow\}^d \mid s \in S\}$$

that is to say that we copy every state of S and encode in them the current d directions of the d heads.⁸ At every transition the d th head moves according to the direction encoded in the state. For $k < i < k'$, when the i th head reaches an endmarker, it changes its direction, something that is encoded in the state. If this endmarker was $\triangleright$, the $i + 1$ th head moves of one square according to its direction. If the k' th head has made a roundtrip on n —that is, is back on $\triangleright$ — M' halts rejecting.⁹

If M did not stop after $\text{Card}(S) \times 2^k \times (n+2)^k$ transitions, it means that it was stuck in a loop and therefore will never accept. By construction, M' always halts after n^d transitions and accepts exactly as M does, so we proved that $\mathcal{L}(M) = \mathcal{L}(M')$ and that M' always halts. **Q.E.D.**

CLAIM 5: *Finite Automata can compute addition, subtraction, power, multiplication, division¹⁰ and modulo over integers.*

Proof. We prove that given two heads H_p and H_q such that $\#H_p = p$ and $\#H_q = q$, $M \in \text{FA}(k)$ with $k \geq 4$ can put a head at distance $q + p$, $q - p$, q^p , $q \times p$ or $\lfloor q/p \rfloor$.

Suppose H_1 and H_2 are at $\triangleright$, they will be used to compute the value we are looking for. We define the following routines:

q + p H_1 goes right until $\#H_1 = \#H_q$, and every time H_1 goes right, H_p moves one square right.

When H_1 reaches H_q , H_p is at distance $q + p$ from $\triangleright$.

q - p H_p goes left until it reaches $\triangleright$, and every time H_p goes left, H_q moves one square left.

When H_p reaches $\triangleright$, H_q is at distance $q - p$ from $\triangleright$. If H_q reaches $\triangleright$ before H_p , p was greater than q .

⁸So if we had $\text{Card}(S) = s$, $\text{Card}(S') = s \times 2^d$, and we should even double the number of states to handle the fact that the $d + 1$ th head has already read $\triangleright$ or not.

⁹That can be simply done by halting the computation in a state not belonging to F .

¹⁰Rounded by the floor operation $\lfloor \cdot \rfloor$, with $\lfloor x \rfloor = \max\{m \in \mathbb{Z} \mid m \leq x\}$.

q^p H_1 goes right until $\#H_1 = \#H_q$

H_1 goes back to $\triangleright$ and every time H_1 moves left, H_q moves one square right.

When H_1 is at $\triangleright$, H_p goes one square left.

We iterate this routine until H_p read $\triangleright$, when it does, $\#H_q = q^p$.

q $\times$ p H_1 goes right until $\#H_1 = \#H_q$

H_2 goes right until $\#H_2 = \#H_q$, and every time H_2 moves, H_p goes one square right.

When H_2 reaches H_q , H_1 goes one square left.

H_2 goes left until it reaches $\triangleright$, every time H_2 moves, H_p goes one square right.

When H_2 reaches $\triangleright$, H_1 goes one square left.

When H_1 reaches $\triangleright$, H_p is at distance $q \times p$ from $\triangleright$.

$\lfloor q/p \rfloor$ H_1 goes right until $\#H_1 = \#H_p$, and every time H_1 moves, H_q goes one square left.

When $\#H_1 = \#H_p$, H_2 moves one square right.

Then H_1 goes left until it reaches $\triangleright$, and every time H_1 moves, H_q goes one square left.

When H_1 reaches $\triangleright$, H_2 moves one square right.

We iterate this routine until H_q reaches $\triangleright$. When it does, H_2 is at distance $\lfloor q/p \rfloor$ from $\triangleright$.

Regarding $q \bmod p$, simply remark that $q \bmod p = q - (\lfloor q/p \rfloor \times p)$, and we saw we can perform any of those operations, provided we have three extra heads to perform some intermediate computation.

About k and time: in the computations as developed above, up to three additional heads are required and the FA is supposed to have sensing heads.¹¹ The routines are linear in q , and they sometimes “destroy” the value q or p , but we may simply add an additional head that will be used as a marker to store the value. **Q.E.D.**

Classical results

To motivate the use of FA, we begin by giving some separation results, and then we study the links between $\mathbf{L}$, $\mathbf{NL}$ and FA.

Theorem 4.2.1 (Separation results). *For all $k \geq 1$,*

$$\mathcal{L}(\text{1DFA}(k)) \subsetneq \mathcal{L}(\text{1DFA}(k+1)) \quad (4.3)$$

$$\mathcal{L}(\text{1NFA}(k)) \subsetneq \mathcal{L}(\text{1NFA}(k+1)) \quad (4.4)$$

$$\mathcal{L}(\text{1DFA}(k)) \subsetneq \mathcal{L}(\text{1NFA}(2)) \quad (4.5)$$

$$\mathcal{L}(\text{DFA}(k)) \subsetneq \mathcal{L}(\text{DFA}(k+1)) \quad (4.6)$$

$$\mathcal{L}(\text{NFA}(k)) \subsetneq \mathcal{L}(\text{NFA}(k+1)) \quad (4.7)$$

$$\mathbf{NL} = \mathbf{L} \text{ iff } \mathcal{L}(\text{1NFA}(2)) \subseteq \bigcup_{k \geq 1} \mathcal{L}(\text{DFA}(k)) \quad (4.8)$$

¹¹But we saw how to get rid of that feature with an additional head previously.

YAO and RIVEST [137] proved eqs. (4.3) to (4.5), eqs. (4.6) and (4.7) are due to MONIEN [108], and the proof of eq. (4.8) may be found in SUDBOROUGH [125, p. 75].

It is often taken for granted that the computation of FA is equivalent to log-space computation on a Turing Machine. For instance, for MONIEN [108, p. 67], “[i]t is wellknown that $\text{SPACE}(\log n)$, the class of languages acceptable within space bound $\log n$, is identical with the class of languages acceptable by two-way multihead automata”. HOLZER, KUTRIB, and MALCHER [75] gives as a reference for this theorem HARTMANIS [69], who introduced those “simple computing devices”. It is true that this paper proves that “log n -tape bounded T.M.’s and k -head automata” are equivalent, but it is only a part of a more complex proof. So we begin with a reformulation of the proof that a FA can simulate a log-space bounded Turing Machine.

Theorem 4.2.2. *For all Turing Machine M using a logarithmic amount of space, there exists $k' \in \mathbb{N}^*$ and $M' \in \text{FA}(k')$ such that $\mathcal{L}(M) = \mathcal{L}(M')$.*

Proof. We can assume without loss of generality that M has k working tapes, that its alphabet is $\{0, 1\}$, that during the course of the computation, M won’t use more than $\log(n)$ cells on each tape, and that all of its cells contains at the beginning of the computation the symbol 0. We will prove that M' accepts the same language as M , with $(k \times 2) + 5$ heads, $H_b^1, H_a^1, \dots, H_b^k, H_a^k, H_r$, and four additional heads to perform some intermediate computation without destroying the values, following Claim 5.¹²

Informally, H_r will act as the reading head of M , H_b^i will encode the bits at the left (“before”) of the head on the i th tape, and H_a^i will encode the bits at the right (“after”) of this head. The bit currently read is encoded, by convention, in H_b^i .

For instance, suppose the head on the i th tape is on the j th cell, and the first j bits on this tape corresponds to the binary writing of the integer e . Then H_b^i will be at distance e from $\triangleright$ (i.e. $\#H_b^i = e$). The remaining $j - \log(n)$ bits are encoded *backward* by an integer e' , and H_a^i will be at distance e' from $\triangleright$ (i.e. $\#H_a^i = e'$). The trick is that the least significant bit is in both case near the head of M' .

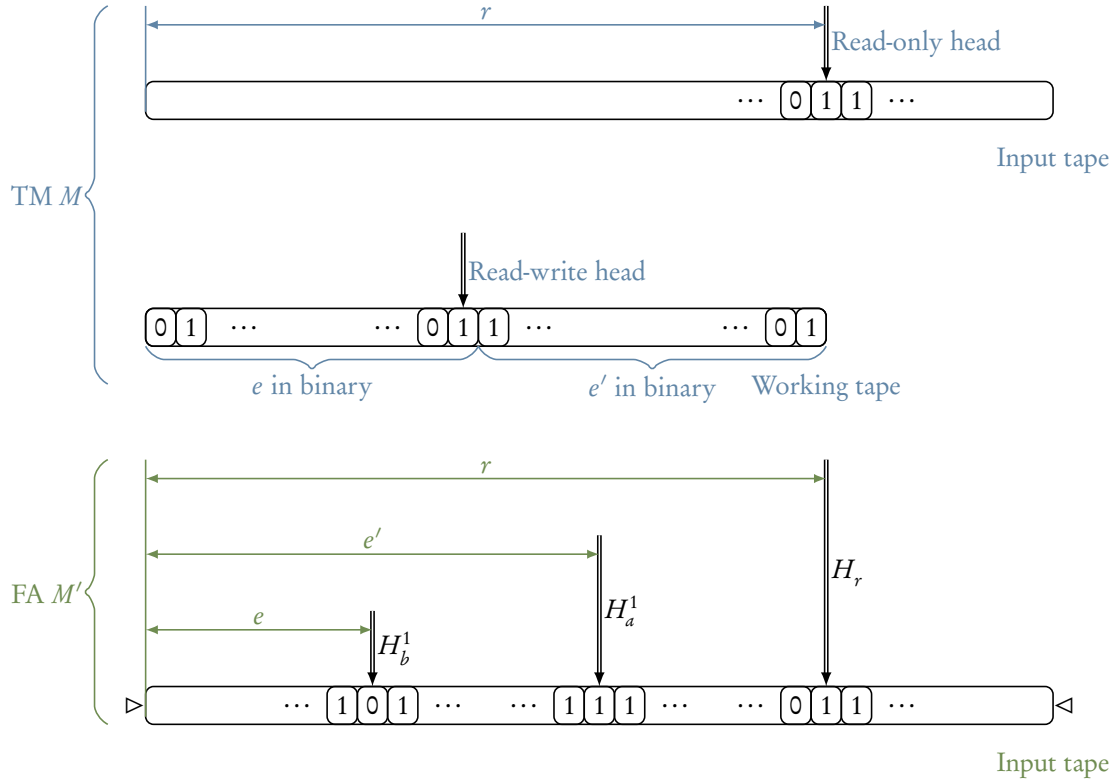
The following table explains how to mimic the basic operation of M with M' :

Turing Machine	Finite Automata
Move reading head	Moves H_r accordingly
Read the value under the head on the i th tape	Compute $\#H_b^i \bmod 2$
Write a 0 instead of a 1 (resp. a 1 instead of a 0) under the head on the i th tape	Move H_b^i one square right (resp. move H_b^i one square left)
Move the head on the i th tape right	Let $\#H_b^i$ be $(2 \times \#H_b^i) + (\#H_a^i \bmod 2)$ Let $\#H_a^i$ be $\#H_b^i / 2$
Move the head on the i th tape left	Let $\#H_a^i$ be $(2 \times \#H_a^i) + (\#H_b^i \bmod 2)$ Let $\#H_b^i$ be $\#H_a^i / 2$

All those intermediate computation may be written as routines, they all take a time linear in $|n|$, even without sensing heads. Those are not states, just the basic of the Turing Machine computation, how it actually runs. By just replacing the reading / writing pattern by movements of heads, we can encode the content of the tapes in the position of the heads.

Since M cannot write more than $\log(n)$ bits on each tape, it is obvious that the heads of M' does not need to be at a distance greater than $|n|$. Moreover, it is quite obvious that $\mathcal{L}(M) = \mathcal{L}(M')$, and that if M is (non-)deterministic, so is M' .

¹²A lot of optimization could be performed on this algorithm, of course, but this is not what is at stake here.



For simplicity, we considered a log-space bounded Turing Machine M with only one working tape. Remember that e' corresponds to the string at the right of the read-write head *encoded backward*, to lighten the simulation of the writing. Of course, the working-tape should be of length logarithmic in the length of the input tape: the scale is not correct!

The Finite Automata M' encodes all the informations within the distance between the beginning of the input tape and the heads. We refer to [Theorem 4.2.2](#) for a precise statement.

Figure 4.2: How to simulate Turing Machines with Finite Automata

Once we know the number of tapes and if the Turing Machine is deterministic or not, it is easy to build a Boolean circuit family that takes as input $(Q, \Sigma, \delta, q_0, g)$ the description of M and that output $\{S, A, k, \triangleright, \triangleleft, s_0, F, \sigma\}$ the description of M' .¹³ This can be done with constant-depth circuit. **Q.E.D.**

The [Figure 4.2](#) should help the reader to get how this simulation works.

Corollary 4.2.3.

$$L \subseteq \text{DFA}$$

$$NL \subseteq \text{NFA}$$

To obtain the converse inclusion is quite trivial: a Turing Machine simulates a $FA(k)$ with $k \times \log(n)$ space by writing and updating the addresses of the k pointers, and that obviously take a logarithmic

¹³That would require to complete the proof to take care of the “special states” of M **accept** and **reject**, and to encode Q and δ in binary, but this is just a matter of plain encoding.

amount of space in n to write an address on a string of size n . If the FA was deterministic, so is the Turing Machine, otherwise they are both non-deterministic.

Theorem 4.2.4.

$$\mathbf{L} = \mathbf{DFA}$$

$$\mathbf{NL} = \mathbf{NFA}$$

What is the exact meaning of this theorem? It is of interest to us because it gives an implicit characterization of both **L** and **NL**: any program that can be written with this device will use a logarithmic amount of space to be simulated on a Turing Machine, any language that can be recognized belong to **L** or **NL**. It is also the reason why 2-ways FA are more studied than their 1-way variant.

It might not be evident at first glance, but Finite Automata in general and this model in particular shares a lot with descriptonal complexity, as illustrated by an (other) excellent survey recently published by HOLZER and KUTRIB [74].

4.3 Non-Deterministic Pointer Machines

All those previous modifications and theorems are classic, but now we will introduce some *ad-hoc* modifications to build a machine model called Pointer Machines (PM). This device is suited to be simulated by operators, it will help us to grasp the computational behaviour of observations. It was introduced and refined in two previous works [10, 11] led with Thomas SEILLER, but the proofs and explanations below are far more advanced than what was proposed.

The main feature of this model is that it is universally non-deterministic, something that has been foreseen by GIRARD:

“[One of the goal of the GoI program is] to define an abstract computer that should not be too far from concrete ones. As we explained in [49], the geometrical analysis should essentially free us from ad hoc time, i.e. artificial causalities, overcontrol, and this means that we are clearly seeking a parallel computer.”

(GIRARD [48, p. 222])

The universally non-deterministic setting is clearly what we are looking for¹⁴: there is strictly no sharing, no communication, in this setting, every computation goes wildly in its own direction, even if the acceptance of the overall process (the nilpotency) is obtained by the cancellation of all the branches.

Pointers Machines (PMs) are designed to mimic the computational behaviour of operators: they do not have the ability to write, their input tape is cyclic, and rejection is meaningful whereas acceptance is the default behaviour. In a sense, the PMs are somehow the dual of FA, concerning acceptance, for a PM rejects if one of its branch reaches **reject**. We will present and motivate the other peculiarities of this model before defining them, and then we will prove that the similarities it shares with FA are “stronger” than their differences. An alternate proof of the characterization of **co-NL** by Non-Deterministic Pointer Machines may be found in [Appendix B](#). This proof is somehow more direct for it does not need FA, but it cannot be extended to deterministic computation and does not provide the elements for a better understanding that are exposed below.

¹⁴And yet we did not use the “Alternating Multihead Finite Automata” proposed by KING [88], for we would have used not enough of their features to justify their introduction.

Presentation

First, we list some peculiarities of any PM:

1. Its alphabet is $\{0, 1, \star\}$ and $\star$ occurs only once.
2. It moves at most one pointer at every transition.
3. It does not have sensing heads.
4. It always stops.¹⁵
5. Its input is written on a circular input tape.
6. Acceptance and rejection are not states nor “labeled states”, but in the co-domain of the transition function. They are, so to say, configurations.
7. Its transition relation is “total”, i.e. for every configuration that is not **accept** or **reject**, there is a transition that can be applied.
8. The initialisation is quite different, because the value is read only when the pointer moves.
9. It suffices that one branch reaches **reject** to make the computation rejects.

We already showed that items 1 to 4 are reasonable modifications of FA, i.e. that provided we are not concerned by the number of heads, we can add or remove those features. We prove in the following (Claim 6, Claim 7) that item 5 and item 7 are not really important,¹⁶ i.e. that we can for free add or remove this feature to FA.

The introduction of another name is justified¹⁷ by item 7 and item 8: they would have needed to strongly modify the definition of FA to be implemented. If we took into account that $\text{co-NL} = \text{NL}$, item 9 could be implemented in the FA. We do not admit this theorem here, for it would mask one of the specificity of our machine, that justifies to see it as a complementary of FA and ease the gateways with operators.

We immediately prove that item 5 is not a big change:

CLAIM 6: *The input tape of a FA can be circular.*

Proof. Let $M \in \text{FA}(k)$ whose input w is written as $\star w$ on a circular input tape i.e. that M can with two transitions goes from the last (resp. first) to the first (resp. last) bit of w . Suppose moreover that M moves only one head at a time.

To simulate M , we define $M' \in \text{FA}(k)$, with $\triangleright \gg \star w \star \ll \triangleleft$ written on its (flat) input tape. The transition function σ' of M' is the same as the transition function σ of M , except that for all state $s \in S$, for all $1 \leq i \leq k$, we add:

$$\begin{aligned} & (s, (*, \dots, *, \gg, *, \dots, *)) \sigma' (s_i^{\text{end}}, (0, \dots, 0, +1, 0, \dots, 0)) \\ & (s_i^{\text{end}}, (*, \dots, *, [\text{not } \triangleleft], *, \dots, *)) \sigma' (s_i^{\text{end}}, (0, \dots, 0, +1, 0, \dots, 0)) \\ & (s_i^{\text{end}}, (*, \dots, *, \triangleleft, *, \dots, *)) \sigma' (s, (0, \dots, 0, -1, 0, \dots, 0)) \end{aligned}$$

¹⁵A property we called *acyclicity* when applied to NDPMs in our works with Thomas SEILLER.

¹⁶The item 7 may be seen as the reverse operation of *minimization* [76, p. 159, section 4.4.3].

¹⁷Beside, we wanted to keep the original naming we developed with Tomas SEILLER.

$$\begin{aligned}
& (s, (*, \dots, *, \ll, *, \dots, *)) \sigma' (s_i^{\text{beg}}, 0, \dots, 0, -1, 0, \dots, 0) \\
& (s_i^{\text{beg}}, (*, \dots, *, [\text{not } \triangleright], *, \dots, *)) \sigma' (s_i^{\text{beg}}, 0, \dots, 0, -1, 0, \dots, 0) \\
& (s_i^{\text{beg}}, (*, \dots, *, \triangleright, *, \dots, *)) \sigma' (s, 0, \dots, 0, 1, 0, \dots, 0)
\end{aligned}$$

where the instructions depends on, and give instructions to the i th head of M' . Those two routines permits to “freeze” the computation when a head is told to “use the ciclicity of the input tape”. We move it to the other side of the tape and then resume the computation. **Q.E.D.**

CLAIM 7: *We can take F the set of accepting states to be a singleton, and make that every branch of the computation halt in that state or in another predefined state.*

Proof. Let M be a FA that always stop, we defined M' with two additional states **accept** and **reject**. For every state $s \in S$ and every configuration $(s, (a_1, \dots, a_k))$ such that $\sigma(s, (a_1, \dots, a_k)) = \emptyset$, we have in σ'

$$(s, (a_1, \dots, a_k)) \sigma' \begin{cases} (\text{accept}, (0, \dots, 0)) & \text{if } s \in F \\ (\text{reject}, (0, \dots, 0)) & \text{elsewhere} \end{cases} \quad (4.9)$$

Q.E.D.

This is an important claim, for it shows how to get closer to a model of computation similar to the Turing Machine, where the computation ends when we reach **accept** or **reject**. By the way, we can notice that the FA M' we defined respects item 7, i.e. its transition relation is “total”.

A Pointer Machine is given by a set of pointers that can move back and forth on the input tape and read (but not write) the values it contains, together with a set of states. For $1 \leq i \leq p$, given a pointer p_i , only one of three different *instructions* can be performed at each step: p_i+ , i.e. “move one step forward”, p_i- , i.e. “move one step backward” and ϵ_i , i.e. “do not move”. In the following, we let $I_{\{1, \dots, p\}}$ be the set of instructions $\{p_i+, p_i-, \epsilon_i \mid i \in \{1, \dots, p\}\}$. We will denote by $\#p_i$ the distance (clockwise) between $\star$ and the current position of p_i .

Definition 4.3.1. A non-deterministic pointer machine (NPM) with $p \in \mathbb{N}^*$ pointers is a couple $M = \{Q, \rightarrow\}$ where Q is the set of states; we always take the alphabet Σ to be $\{0, 1, \star\}$ and the set of instructions to be $I_{\{1, \dots, p\}}$; and $\rightarrow \subseteq (\Sigma^p \times Q) \times ((I_{\{1, \dots, p\}}^p \times Q) \cup \{\text{accept}, \text{reject}\})$ is the *transition relation* and it is total. We write $\text{NPM}(p)$ the set of non-deterministic pointer machines with p pointers.

We define a *pseudo-configuration* c of $M \in \text{NPM}(p)$ as a “partial snapshot”: $c \in \Sigma^p \times Q$ contains the last values read by the p pointers and the current state, *but does not contain the addresses of the p pointers*. The set of pseudo-configurations of a machine M is written C_M and we should remark that it corresponds to the domain of the transition relation. If $\rightarrow$ is functional, M is a *deterministic* pointer machine, and we write $\text{DPM}(p)$ the set of deterministic pointer machines with p pointers. We let $\text{PM}(p) = \text{DPM}(p) \cup \text{NPM}(p)$.

Let $M \in \text{PM}(p)$, $s \in C_M$ and $n \in \mathbb{N}$ an input. We define $M_s(n)$ as M with n encoded as a string on its circular input tape (as $\star a_1 \dots a_k$ for $a_1 \dots a_k$ the binary encoding of n and $a_{k+1} = a_0 = \star$) starting in the *initial pseudo-configuration* s with $\#p_i = 0$ for all $1 \leq i \leq p$ (that is, the pointers starts on $\star$). Each of them is associated to a *memory slot* (or *register*) that store the values it reads, but (recall

item 8) as the pointers did not moved yet, no value has been read, and the slots are empty. The initial pseudo-configuration s gives the p values for the p registers and the initial state s .

Remark that the *initial pseudo-configuration* s does not initializes those p registers necessarily in a faithful way (it may not reflect the values contained at $\#p_i$). That is to say that the first transition will rely on the values enclosed in the pseudo-configuration, and not on the values the pointer actually points at (that are always, by definition, $\star$). This complex definition will be justified in the following section.

The entry n is *accepted* (resp. *rejected*) by M with *initial pseudo-configuration* $s \in C_M$ if after a finite number of transitions every branch of $M_s(n)$ reaches **accept** (resp. at least a branch of M reaches **reject**). We say that $M_s(n)$ halts if it accepts or rejects n and that M decides a set S if there exists a pseudo-configuration $s \in C_M$ such that $M_s(n)$ accepts if and only if $n \in S$.

We write as usual $\mathcal{L}(M)$ the language of a PM M , and **DPM** (resp. **NPM**) the set of languages decided by $\bigcup_{p \in \mathbb{N}^*} \text{DPM}(p)$ (resp. $\bigcup_{p \in \mathbb{N}^*} \text{NPM}(p)$).

A modification of Finite Automata

One could quite simply prove that PM are just reasonable re-arranging of the presentation of FA “of a special kind”. The minor modifications (among which the initial pseudo-configuration, the fact that **accept** and **reject** are not in the set of the states) are just red tape and does not enhance nor lower the power of computation of FA. What is really at stake, here, is to take into account that **co-NL** = **NL** or not. If we do, then it is obvious that PM can recognise any language in **NL**, if we don’t, we should rather see them as some kind of complementary of FA. We prefer not to take as granted this result for the computation of operator algebra we are going to mimic should really be understood as universally non-deterministic.

Proposition 4.3.1 (Simulation of a FA by a PM). *For all $k \in \mathbb{N}^*$, for all $M \in \text{NFA}(k)$, there exists $M' \in \text{PM}(k+2)$ such that $\text{co-}\mathcal{L}(M) = \mathcal{L}(M')$.*

Proof. We can assume that M always halts, does not have sensing heads and has for alphabet $\{0, 1\}$. The construction of M' is really simple: given $w \in \{0, 1\}^*$ an input, M' takes as input the description of M as a list $S \triangleright w \triangleleft \sigma s_0 F$. Then M' is initialized by placing k pointers on $\triangleright$, a pointer on s_0 and the last one on the first symbol of σ . This last pointer scans σ to find a transition that can be applied. To simulate a transition, M' moves the k first pointers according to σ and the $k+1$ th pointer is moved to a new state according to σ .

When a transition that may be applied is found, a non-deterministic transition takes place: on one way M' simulates this first transition, on the other way M' keeps scanning σ to look for another transition to apply. When no transition may be applied —and it always happens, as M always halts—, M' checks if the current state belongs to S : if it does, M' rejects, elsewhere M' accepts.

It can be easily verified that M' accepts iff M rejects, that $k+2$ pointers are enough to perform this simulation, and that if M is deterministic, so is M' , because there is no need to look for another transition to apply. Note moreover that M' always halts, as M does. **Q.E.D.**

Corollary 4.3.2.

$$\mathbf{L} \subseteq \mathbf{DPM}$$

$$\mathbf{co-NL} \subseteq \mathbf{NPM}$$

We do not prove the reverse inclusion by proving that FA can simulate PM, even if that would not be difficult. Instead, we will prove that the model we defined in the previous chapter, operators, can simulate PM: together with the results of the previous chapter, that gives us the reverse inclusion.

4.4 Simulation of Pointer Machines by Operator Algebra

This section recalls the results previously obtained [10] and the enhancements that were proposed later [11]. We tried to use the space at our disposal to ease the comprehension of this simulation: we will recall the framework we are working with, give a general idea of the simulation, introduce some notations and then encode the basic operations of a PM. Our efforts to define properly the basic mechanisms will give us the main result of this section almost for free, but a last subsection will be devoted to the careful analysis of the deterministic case.

Our aim in this section is to prove (Lemma 4.4.1) that for any PM M and initial pseudo-configuration $s \in C_M$, there exists an observation $M_s^\bullet \in M_6(\mathfrak{G}) \otimes Q_M$ such that for all $N_n \in M_6(\mathfrak{N}_0)$ a binary representation of n , $M_s(n)$ accepts if and only if $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ is nilpotent. We will prove moreover that $M_s^\bullet \in P_+$, and that if M is deterministic, then $M_s^\bullet \in P_{+,1}$, i.e. that the encoding of DPMs satisfies a particular property: their 1-norm is less or equal to 1.

We will proceed carefully, for the encoding is a bit unintuitive, mainly because the computation as simulated by operators is *widely parallel*.

“[O]ur modelisation is without global external time, i.e. without any kind of synchronisation. Due to this lack of synchronisation, some ‘actions at distance’ familiar to the syntax cannot be performed [...]. We have therefore gained a new freedom w.r.t. time, and this should be essential in view of parallel execution.”

(GIRARD [48, p. 223])

This freedom regarding time is for the moment hard to capture, and that is one of the reasons we are going to simulate non-deterministic log-space computation, in a highly parallel way.

Recalling the framework

The setting is as we left it in Chapter 3: we take the binary representation of integers N_n as operators in $M_6(\mathfrak{N}_0)$ (Definition 3.4.1), we consider only the normative pair $(\mathfrak{N}_0, \mathfrak{G})$ as defined in Corollary 3.4.4 and we take the sets of observations $P_{\geq 0}$, P_+ and $P_{+,1}$ defined in Definition 3.4.7.

Intuitively, the encoding of $M = \{Q, \rightarrow\}$ a PM with p pointers and an initial pseudo-configuration s will be an observation $M_s^\bullet$, which is an operator of $M_6(\mathfrak{G}) \otimes Q_M$, where the algebra of states¹⁸ Q_M is

$$Q_M = \underbrace{M_6(\mathbb{C}) \otimes M_6(\mathbb{C}) \otimes \cdots \otimes M_6(\mathbb{C})}_{p \text{ times}} \otimes M_{q'}(\mathbb{C})$$

This algebra Q_M represents a set of *pseudo-states*, i.e. it will encode the memory slots that store the last p values read by the pointers, and the extended set of states $Q^\uparrow$. We extend the set of states Q of M with

¹⁸That should rather be called the *algebra of pseudo-states*, but we wanted to stay consistent with the notation of Chapter 3, and, besides, the state will also be encoded in it.

some intermediate states needed to encode the basic operations of M properly. We let $\text{Card}(Q^\uparrow) = q'$ and will call an *extended pseudo-configuration* an element of $\{0, 1, \star\}^p \times Q^\uparrow$.

The intuition is that the j th copy of $M_6(\mathbb{C})$ represents the memory slot that contains the last value read by the j th pointer. We will therefore distinguish for each copy of $M_6(\mathbb{C})$ a basis $(0o, 0i, 1o, 1i, s, e)$ corresponding to the different values that a pointer can read.¹⁹ To sum up, the distinguished basis of Q_M considered will be denoted by tuples $(a_1, \dots, a_p, \mathbf{q})$. Notice that such a tuple naturally corresponds to an extended pseudo-configuration.

A crucial remark should be made to explain properly how the computation will be performed. Recall that the input of $M_s^\bullet$ will be $N_n \otimes \text{Id}_{Q_M}$, i.e. the tensoring of N_n with the unit of the algebra of states. For the product $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ to be nilpotent means to be nilpotent for *any basis* of Q_M , that is, for any possible initial pseudo-configuration.

Stated differently, the integer is considered at the same time in every possible pseudo-configuration. As a result, the computation represented by the sequence

$$M_s^\bullet(N_n \otimes \text{Id}_{Q_M}), (M_s^\bullet(N_n \otimes \text{Id}_{Q_M}))^2, \dots$$

somehow simulates for all $c \in C_M$ the computations of $M_c(n)$ simultaneously. Of course, this applies only for the initialisation: once the first pseudo-configuration is fixed, the computation is led according to it. But among all the computations started in parallel, only one effectively has for initial pseudo-configuration s .

The representation of reject cannot be done without considering an initial pseudo-configuration, because if the computation starts with the “wrong” pseudo-configuration, we will under certain circumstance start it again with the “good” initial pseudo-configuration, i.e. s . This will be explained when we will encode rejection.

Encoding a machine

In fact, our operator could be conceived as having a single pointer, and p registers. This explains the “one movement at a time” motto, our single pointer will travel a lot, but this is algorithmically equivalent to having p pointers. GIRARD [59, p. 262] refers to this oddity as if we had at our disposal a “thumb” and several fingers. Only the thumb can move, but it can exchange its position with any other finger.

We know how to access the values of N_n thanks to projections, but we have to handle two difficulties: that every other bit is redundant, and that the instruction will depend on the previous projection, “go deeper” or “lift back to the surface”. What we mean is that changing the direction of the movement will not be trivial: we will handle differently a movement from left to right on the string (a “forward move”) and a movement from right to left (a “backward move”), and furthermore we will handle them differently depending on the previous orientation of the movement. Decomposing this movement and handling the different case will require some intermediate steps and *extended states*.

We will encode the transition function²⁰ of M as if it were a function between C_M and $C_M \cup \{\text{accept}, \text{reject}\}$. This could seem really surprising, for establishing the pseudo-configuration c' reached

¹⁹As Thomas SEILLER points it, we should rather consider those values dynamically: our observation will always be on the edge of reading the next value, so reading an occurrence of $0o$ rather corresponds to reading the input of the value connected to this output of a 0 .

²⁰We continue for the sake of clarity to speak of $\rightarrow$ as if it was a function, but it is of course eventually a relation, in the non-deterministic case.

after applying the transition $\rightarrow$ from a pseudo-configuration c needs to take care of the actual input and of the positions of the pointers.²¹ Here the actual position of the pointers and the reading of the input will be handled differently: the actual operation of reading the input and moving the pointers will be handled directly within the transition, i.e. c' will be reached iff the values read correspond to a “legal move” from c to c' . So to say, we will skip the “instruction part” to go directly to the pseudo-configuration resulting from the movement of *the* pointer and the reading of *the* new value.²²

So we will encode each couple (c, c') of pseudo-configuration such that “nothing in $\rightarrow$ prevents c' to be reached from c ” by an operator $\phi_{c, c'}$. How can we know that $c \rightarrow c'$ is possible without knowing nor the input nor the position of the heads? Well, every time there is $c \rightarrow (i_1, \dots, i_p, \mathbf{q}')$ for i instructions $i_1, \dots, i_p \in I_{\{1, \dots, p\}}$, we consider that every pseudo-configuration c' whose state is $\mathbf{q}'$ can be reached from c . The following will explain how those pseudo-configurations are sorted between the one that are actually reachable considered an input and the one that are not.

All the possibles transitions from the pseudo-configuration c are then encoded by

$$\sum_{c \rightarrow c'} \phi_{c, c'}$$

and the encoding of the transition relation will then correspond to the sum:

$$M_s^\bullet = \sum_{c \in C_M} \sum_{c' \text{ s.t. } c \rightarrow c'} \phi_{c, c'} + \text{reject}_s \quad (4.10)$$

As we may see, $M_s^\bullet$ is entirely defined by the encoding of the transition function of M , decomposed in all its possible applications. Every transition from every pseudo-configuration will be encoded in the sum, and a built-in way of assuring that it corresponds to the actual position of the pointers and the input will discriminate between the transitions that are actually possible and the one that aren't.

In fact the instructions to move a pointer will be performed by the selection of the pointer, a projection onto the values that may have been read, the storing of the value read, the deactivation of the pointer and the changing of state.

Notations

Before describing the encoding in details, we need to introduce several notations that will be used for the definition of the $\phi_{c, c'}$ operators. The von Neumann algebra $M_6(\mathfrak{G}) \otimes Q_M$ will be considered as $M_6(\mathbb{C}) \otimes \mathfrak{G} \otimes Q_M$ and we will define the operators needed to encode the basic operation of M as tensor products $u \otimes v \otimes w$, where $u \in M_6(\mathbb{C})$, $v \in \mathfrak{G} \subseteq \mathfrak{R}$ and $w \in Q_M$. We will explain the main patterns of those three operators.

Some of these definitions are tiresome, because the values of N_n are encoded in the pattern we defined in Chapter 3, and so every bit comes in a “input” and “output” version. That is to say, if we remember Figure 3.4 and more generally Section 3.3, that $\star a_k \dots a_1 \star$ is in fact encoded as $e a_{i_1} a_{o_1} \dots a_{i_k} a_{o_k} s$. Moreover, we need to cope²³ with the “last direction of the pointer” to adapt cleverly the projections.

The operator $u \in M_6(\mathbb{C})$ will alternatively read a value, and enforce the direction. We define the projections $\pi_\diamond$ of $M_6(\mathbb{C})$ for $\diamond \in \{0i, 0o, 1i, 1o, e, s\}$ as the projections on the subspace induced by the

²¹Or, equivalently, of the position of our “traveler” and of the values of the registers.

²²For M move only one pointer at a time.

²³This comes from a concern regarding the norm of the operators we are going to define, something we did not bother with in our first work [10, p. 18]. In this work, [in] and [out] directly encodes the movement, regardless of the orientation of the previous movement.

basis element $\blacklozenge$. We will sometimes denote e (resp. s) by $\star i$ (resp. $\star o$). The reading part will be essentially handled by those projections.

The direction is managed by

1. $\pi_{\text{in}} = \pi_{0i} + \pi_{1i} + \pi_e$ and $\pi_{\text{out}} = \pi_{0o} + \pi_{1o} + \pi_s$. They are such that $\pi_{\text{in}}\pi_{\text{out}} = \pi_{\text{out}}\pi_{\text{in}} = 0$.
2. $[\text{in} \rightarrow \text{out}]$ and $[\text{out} \rightarrow \text{in}]$ are such that

$$\begin{aligned} [\text{in} \rightarrow \text{out}][\text{out} \rightarrow \text{in}] &= \pi_{\text{out}} & [\text{out} \rightarrow \text{in}][\text{in} \rightarrow \text{out}] &= \pi_{\text{in}} \\ [\text{out} \rightarrow \text{in}]\pi_{\text{out}} &= [\text{out} \rightarrow \text{in}] & [\text{in} \rightarrow \text{out}]\pi_{\text{in}} &= [\text{in} \rightarrow \text{out}] \\ [\text{out} \rightarrow \text{in}]\pi_{\text{out}} &= 0 & [\text{in} \rightarrow \text{out}]\pi_{\text{in}} &= 0 \end{aligned}$$

They are “the memory of the direction of the previous movement”, but one should notice that they are not projections, i.e. $[\text{out} \rightarrow \text{in}]^2 = [\text{in} \rightarrow \text{out}]^2 = 0$.

We can express those operators as matrices, to ease the checking of the previous equalities:

$$\begin{aligned} \pi_{\text{in}} &\equiv \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} & [\text{in} \rightarrow \text{out}] &\equiv \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \end{pmatrix} \\ \pi_{\text{out}} &\equiv \begin{pmatrix} 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix} & [\text{out} \rightarrow \text{in}] &\equiv \begin{pmatrix} 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 0 \end{pmatrix} \end{aligned}$$

The operator $v \in \mathfrak{G} \subset \mathfrak{R}$ will in fact always be a permutation. We will only need in the following $\tau_{i,j}$, the transposition exchanging the integers i and j .²⁴

The last operator $w \in Q_M$ will encode an updating in the registers and a changing of state, i.e. the evolution from an extended pseudo-configuration $c = (a_1, \dots, a_p, \mathbf{q})$ to an extended pseudo-configuration $c' = (a'_1, \dots, a'_p, \mathbf{q}')$. We define the partial isometry:

$$(c \rightarrow c') = (a_1 \rightarrow a'_1) \otimes \dots \otimes (a_p \rightarrow a'_p) \otimes (\mathbf{q} \rightarrow \mathbf{q}')$$

where $(v \rightarrow v')$, for $v, v' \in \{0i, 0o, 1i, 1o, e, s\}$ or $v, v' \in Q^\uparrow$, is simply defined as:

$$(v \rightarrow v') := \begin{pmatrix} & & v & & \\ 0 & \dots & 0 & \dots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & 1 & \dots & 0 \\ \vdots & \ddots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & \dots & 0 \end{pmatrix} v'$$

²⁴To be more precise, $\tau_{i,j}$ is the unitary induced by the transposition $\begin{pmatrix} i & j \end{pmatrix}$.

We should remark that for all $c, c', c'' \in C_M$,

$$(c \rightarrow c')c'' = \begin{cases} c' & \text{if } c' = c'' \\ 0 & \text{elsewhere} \end{cases} \quad (4.II)$$

A transition will always impact on at most one value stored in the registers and on the state $\mathbf{q}$. Suppose it acts on the j th register, we define

$$([a_j \rightarrow a'_j]_j, \mathbf{q} \rightarrow \mathbf{q}')$$

as

$$(\text{Id}, \dots, (a_j \rightarrow a'_j), \dots, \text{Id}, \mathbf{q} \rightarrow \mathbf{q}')$$

It is of course possible to encode a transition that let the state unchanged: just let $\mathbf{q} \rightarrow \mathbf{q}'$ be Id if $\mathbf{q} = \mathbf{q}'$.

In practise, we won't need to be specific about a_j and a'_j at the same time. All that will matter will be the value stored *or* the value we are going to store in the j th register, for we will decompose the action of reading the value and the action of storing it in two steps.²⁵

We will use two shorthands to encode the actions on the j th register. For $* \in \{0, 1, \star\}$ and $d \in \{o, i\}$, we define

$$\begin{aligned} [\rightarrow *d]_j &= \sum_{\diamond \in \{0i, 0o, 1i, 1o, e, s\}} [\diamond \rightarrow *d]_j \\ [*d \rightarrow]_j &= \sum_{\diamond \in \{0i, 0o, 1i, 1o, e, s\}} [*d \rightarrow \diamond]_j \end{aligned}$$

At last, we will sometimes let the $*$ of the previous equation be a *subset* of $\{0, 1, \star\}$.

This somehow complicated way of writing the operators could be explained with two examples:

$[\{0i, 1i, \star i\} \rightarrow]_j$ should be read as “no matter what the previous value stored in the j th slot was, we just make sure that it came from a forward move, i.e. it is endowed with a i ”.

$[\rightarrow 0i]_j$ will simply be “store a $0i$ in the j th slot, no matter what the previous stored value was”.

The left hand side of the arrow tells us something about the value a_j of c , the right hand side of the arrow tells us something about the value a'_j of c' .

We are now ready to define the operators needed to encode the basic operations of the PM. acceptance will be especially easy, but basic movement and rejections are a more complex sum of operators defined in Figure 4.3 and Figure 4.4.

Basic operations

From now on, we consider given M a PM with p pointers and $c = (a_1, \dots, a_p, \mathbf{q}) \in C_M$ a pseudo-configuration. An initial pseudo-configuration $s \in C_M$ will be needed to encode rejection. There are only three cases in M : either $c \rightarrow c'$ for c' a pseudo-configuration (but remark that c' is not unique if M is non-deterministic), $c \rightarrow \text{accept}$, or $c \rightarrow \text{reject}$. We will define the encoding relatively to those three cases.

$$\begin{aligned}
bf_{j,q}^c &= \pi_{out} \otimes \tau_{0,j} \otimes ([\{0o, 1o, \star o\} \rightarrow]_j, q \rightarrow \mathbf{mov}_j^{q,q'}) \\
ff_{j,q}^c &= [\text{in} \rightarrow \text{out}] \otimes \tau_{0,j} \otimes ([\{0i, 1i, \star i\} \rightarrow]_j, q \rightarrow \mathbf{mov}_j^{q,q'}) \\
fb_{j,q}^c &= \pi_{in} \otimes \tau_{0,j} \otimes ([\{0i, 1i, \star i\} \rightarrow]_j, q \rightarrow \mathbf{mov}_j^{q,q'}) \\
bb_{j,q}^c &= [\text{out} \rightarrow \text{in}] \otimes \tau_{0,j} \otimes ([\{0o, 1o, \star o\} \rightarrow]_j, q \rightarrow \mathbf{mov}_j^{q,q'}) \\
frec_{j,q}^c &= \sum_{* \in \{0,1,\star\}} (\pi_{*i} \otimes \tau_{0,j} \otimes ([\rightarrow *i]_j, \mathbf{mov}_j^{q,q'} \rightarrow q')) \\
brec_{j,q}^c &= \sum_{* \in \{0,1,\star\}} (\pi_{*o} \otimes \tau_{0,j} \otimes ([\rightarrow *o]_j, \mathbf{mov}_j^{q,q'} \rightarrow q'))
\end{aligned}$$

$bf_{j,q}^c$ encodes “if the previous movement was a backward move, change the direction, activate the j th pointer, erase the j th register”.

$ff_{j,q}^c$ encodes the same operation in the case where the previous movement was already a forward move.

$fb_{j,q}^c$ and $bb_{j,q}^c$ encodes the same couple of operations for a backward movement.

$frec_{j,q}^c$ is made of a sum over $* \in \{0,1,\star\}$, and every of its member encode “put the j th pointer back to its position, store $*$ in the j th register and adopt state q' .” $brec_{j,q}^c$ encodes the same operation for the other direction.

Figure 4.3: The operators encoding the forward and backward move instructions

Moving a pointer, reading a value and changing the state Suppose that M , in pseudo-configuration $c = (a_1, \dots, a_p, q)$, moves the j th pointer right, reads the value a'_j stored at $\#p_j$, updates the j th memory slot in consequence and changes the state to q' . The new pseudo-configuration is then $c' = (a_1, \dots, a_{j-1}, a'_j, a_{j+1}, \dots, a_p, q')$.

There are two cases: either the last movement of the j th pointer was a forward move or it was a backward move. The case we are dealing with is obtained from the value stored in the j th memory slot: if the value is $0i$, $1i$ or $\star i$, then the last move was a forward move, if the value is $0o$, $1o$ or $\star o$, the last move was a backward move. In the first case, the operator $ff_{j,q}^c$ will be applied, and in the second the $bf_{j,q}^c$ operator will be applied. Remark that we have to make the sum of those two operators, for we cannot foresee what was the direction of the last movement. The fact that the “wrong case” gets to 0 should be evident from the definitions of $[\text{in} \rightarrow \text{out}]$, $[\text{out} \rightarrow \text{in}]$, π_{in} and π_{out} .

Both these operators somehow *activate* the j th pointer by using the transposition $\tau_{0,j}$ and prepare the reading of the representation of the integer. This representation will then give the value of the next (when moving forward) digit of the input. The $frec_{j,q}^c$ or the $brec_{j,q}^c$ operator is then applied in order to simultaneously update the value of the j th memory slot and *deactivate* the j th pointer. Remark that this careful decomposition requires the introduction of a new (i.e. extended) state $\mathbf{mov}_j^{q,q'}$.

To sum up, we let

$$\text{forward}_{j,q}^c = bf_{j,q}^c + ff_{j,q}^c + frec_{j,q}^c \quad \text{backward}_{j,q}^c = fb_{j,q}^c + bb_{j,q}^c + brec_{j,q}^c$$

We insist one last time: this sum is here to cope with all the possible situations regarding the orientation of the previous movement of the j th pointer and the value actually read. All the “incorrect situations” will get cancelled and only the “legal transition between two pseudo-configurations” will be applied.

In the following, we will forget about the subscripts and superscripts of these operators and write ff , fb , bf , bb , and rec for $brec$ or $frec$.

²⁵That will explain the introduction of additional states to handle this intermediate step.

This —by the way— justifies once again the “one movement at every transition” (item 2), for we apply only one permutation at every transition.

Accept The case of acceptance is especially easy: we want to stop the computation, so every transition $(a_1, \dots, a_n, \mathbf{q}) \rightarrow \text{accept}$ will be encoded by 0.

This justifies the item 7 of the previous section: we need the transition function to be always defined. If there is a pseudo-configuration h such that $h \rightarrow \emptyset$, this will be encoded by the operators as $h \rightarrow 0$, hence the default behaviour is to accept when the function is not defined.

So if the computation of M does not lead to acceptance, it is because M rejects (thanks to the “always stops” item 4). On the other hand, if the computation of $M_s^\bullet$ halts, it is in any case because it accepts, so we want to mimic the rejection by a loop. An attempt to define a naive “loop projection” fails: to illustrate it, we define an observation $\text{reject}_{\text{naive}} \in M_6(\mathbb{C}) \otimes \mathfrak{G} \otimes Q_M$ that does nothing but reject, thanks to a yet-to-be-defined “loop projection” π_{reject} :

$$\text{reject}_{\text{naive}} = \text{Id}_{M_6(\mathbb{C})} \otimes \text{Id}_{\mathfrak{R}} \otimes \pi_{\text{reject}}$$

It succeeds to make the computation loops, as for all $d \in \mathbb{N}$ and N_n :

$$\begin{aligned} ((N_n \otimes \text{Id}_{Q_M}) \text{reject}_{\text{naive}})^d &= ((N_n \otimes \text{Id}_{\mathfrak{R}} \otimes \text{Id}_{Q_M}) \text{Id}_{M_6(\mathbb{C})} \otimes \text{Id}_{\mathfrak{R}} \otimes \pi_{\text{reject}})^d \\ &= (N_n \otimes \pi_{\text{reject}})^d \\ &= N_n^d \otimes \pi_{\text{reject}} \end{aligned}$$

And we know that N_n is not nilpotent, so neither is $(N_n \otimes \text{Id}_{\mathfrak{R}}) \text{reject}_{\text{naive}}$: this attempt succeeds in creating a loop.

On the other hand, as the encoding of $\rightarrow$ is built as a sum of the basic operations, π_{reject} appears in it, and $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ cannot be nilpotent. So this solution is excessive and we have to find another way to be sure that the operator will loop if *and only if* the operator that simulates the reject is reached. To do so, we build an operator that simulates the reinitialisation of M in pseudo-configuration s with its p pointers on $\star$ when reached.

It comes with another advantage: remember that the computation of $M_s^\bullet$ started in all possible configurations, not only in s , so in all the other cases, if reject was reached after starting the computation with a pseudo-configuration $c' \neq s$, we enforce the computation of the machine on s . As a consequence, if $M_s(n)$ accepts, the rejection of $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ that corresponds to a computation on c' will be temporary: once rejection attained, the computation restarts with pseudo-configuration s and will therefore halts accepting. The fact that $c' \neq s$ could lead to acceptance is not relevant: as we said (item 9), “acceptance is the default behaviour”, what really matters is rejection.

Reject Rejection should lead to the creation of a loop, in fact it is encoded by “put the p pointers back to $\star$ and starts the computation again with initial pseudo-configuration s ”. This is why we needed the initial pseudo-configuration s to be a parameter, for the reject_s operator will depend of it. We also need two extended states, **back_i** and **mov-back_i** for each $i = 1, \dots, p$.

The transitions of the form $(a_1, \dots, a_n, \mathbf{q}) \rightarrow \text{reject}$ are encoded by the operator that represents the “transition” $(a_1, \dots, a_n, \mathbf{q}) \rightarrow (a_1, \dots, a_n, \text{back}_1)$.

We can then define

$$\text{reject}_s = \sum_{i=1}^p (\text{rm}_i + \text{rr}_i^0 + \text{rr}_i^1 + \text{rc}_i)$$

Remember we are given an initial pseudo-configuration $s = (a_1, \dots, a_n, \mathbf{q}_0)$.

$$\begin{aligned} \text{rm}_i &= \text{Id} \otimes \tau_{0,i} \otimes (\mathbf{back}_i \rightarrow \mathbf{mov-back}_i) \\ \text{rr}_i^0 &= \pi_{0o} \otimes \tau_{0,i} \otimes ([\rightarrow 0o]_i, \mathbf{mov-back}_i \rightarrow \mathbf{back}_i) \\ \text{rr}_i^1 &= \pi_{1o} \otimes \tau_{0,i} \otimes ([\rightarrow 1o]_i, \mathbf{mov-back}_i \rightarrow \mathbf{back}_i) \\ \text{rc}_i &= \begin{cases} \pi_\star \otimes \tau_{0,i} \otimes ([\rightarrow a_i]_i, \mathbf{mov-back}_i \rightarrow \mathbf{back}_{i+1}) & (1 \leq i < p) \\ \pi_\star \otimes \tau_{0,p} \otimes ([\rightarrow a_p]_p, \mathbf{mov-back}_p \rightarrow \mathbf{q}_0) & (i = p) \end{cases} \end{aligned}$$

rm_i encodes “if you are in state $\mathbf{back}_i$, no matter what you read, activate the i th pointer and go to state $\mathbf{mov-back}_i$ ”. rr_i^0 (resp. rr_i^1) encodes “if you read a 0 (resp. a 1), deactivates the i th pointer, store the value you read and go to state $\mathbf{back}_i$ ”.

rc_i encodes “If you read $\star$, deactivates the i th pointer, put the i th value of the initial pseudo-configuration s in the i th register, and change state to $\mathbf{mov-back}_{i+1}$. If you were dealing with the last pointer, go to the state $\mathbf{q}_0$ provided by the initial pseudo-configuration.”

Figure 4.4: The operators encoding rejection

How does this operator encodes a reinitialization of M ? The operator rm_i initializes the process, rr_i^0 and rr_i^1 encodes “as long as the i th pointer is reading a 0 or a 1, keep moving it backward”. We are forced to alternate between the extended states $\mathbf{back}_i$ and $\mathbf{mov-back}_i$ to handle the redundancy of the input. When $\star$ is read by the i th pointer, the actual value of the initial pseudo-configuration s is stored in the corresponding register and rc_i starts this process for the $i + 1$ th pointer.

Lemma 4.4.1. *For all $k \in \mathbb{N}$, for all $M \in \text{NPM}(k)$, for all $s \in C_M$ an initial pseudo-configuration, there exists an operator $M_s^\bullet$ such that for all $n \in \mathbb{N}$ and every binary representation $N_n \in \text{M}_6(\mathfrak{N}_0)$ of n , $M_s(n)$ accepts iff $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ is nilpotent.*

Proof. We obtain $M_s^\bullet$ thanks to the encoding we just defined. We can remark that $M_s^\bullet \in P_+$, and as we are acting in a normative pair, it does not matter which representation N_n of n we pick. So let us fix $n \in \mathbb{N}$ and N_n one of its binary representation.

Considering the representation of the rejection it is clear that if a branch of $M_s(n)$ rejects, the operator $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ will not be nilpotent, so we just have to prove that if $M_s(n)$ accepts then $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ is nilpotent.

We prove its reciprocal: suppose $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ is not nilpotent. In this product N_n is given to the operator $M_s^\bullet$ that starts the simulation of the computation of M with input n in every possible initial pseudo-configuration at the same time. Since the encoding of M is built with respect to s , we know that there exists a j such that $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})\pi_j$ is the simulation of $M_s(n)$, but the computation takes place in the other projections too: for $i \neq j$ it is possible that $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})\pi_i$ loops where for a d $(M_s^\bullet(N_n \otimes \text{Id}_{Q_M}))^d \pi_j = 0$.

We can handle this behaviour because for all $c \in C_M$, M_s always halts. This is due to our “built-in clock”, that insures us that M will always stops, no matter what the initial configuration considered is. So if $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ is not nilpotent it is because at some point the **reject** state has been reached. After this state is reached (let’s say after $r \in \mathbb{N}$ iterations) we know that $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})^r \pi_i$ is exactly the simulation of $M_s(n)$. If it loops again, it truly means that $M_s(n)$ rejects.

So we just proved that $M_s^\bullet(N_n \otimes \text{Id}_{Q_M})$ is not nilpotent iff $(M_s^\bullet(N_n \otimes \text{Id}_{Q_M}))^d \pi_j \neq 0$ for all $d \in \mathbb{N}$. But it is clear that in this case $M_s(n)$ rejects. **Q.E.D.**

Proposition 4.4.2.

$$\mathbf{co}\text{-NL} \subseteq \{P_+\} \subseteq \{P_{\geq 0}\}$$

Proof. Even if the previous simulation may be hard to grasp, to output the description of $M_s^\bullet$ form $M = \{\rightarrow, Q\} \in \text{NPM}(p)$ is $\mathbf{AC}^0$: all we have to do is to list all the $3^p \times \text{Card}(Q)$ possible pseudo-configurations of M , and to output $\phi_{c,c'}$ for every of them. The pseudo-configuration c' is simply obtained by looking in $\rightarrow$ if there is a transition from c to a set of instructions (no matter what they are) endowed with the state of c' . Outputting the members of the sum can obviously be done in parallel, for it relies on simple routines. The reject_s operator is simply a constant factor in the processing.

Thus, the previous proposition tells us that $\mathbf{NPM} \subseteq \{P_+\}$ since the representation $M_s^\bullet$ of a couple (M, s) , where M is an NPM and $s \in C_M$, is obviously in P_+ . Moreover, since $P_+ \subseteq P_{\geq 0}$, we have $\{P_+\} \subseteq \{P_{\geq 0}\}$. The Corollary 4.3.2 gives us the last element to conclude, i.e. that $\mathbf{co}\text{-NL} \subseteq \mathbf{NPM}$. **Q.E.D.**

The deterministic case

This subsection is dedicated to the somehow tedious verification that if M is a DPM, then $M_s^\bullet$ will be of 1-norm (Definition 3.4.6) less than 1. This will conclude this chapter by proving that operators can also mimic deterministic computation.

Proposition 4.4.3. *Let A, B be operators such that $AB^* = B^*A = 0$. Then:*

$$\|A + B\| = \max\{\|A\|, \|B\|\}$$

Proof. Actually, this proposition is proved by MAHER [99, Theorem 1.7 (b)], but stated differently, with conditions on ranges.²⁶ It is proven that if $\text{Ran}(A) \perp \text{Ran}(B)$ and $\text{Ran}(A^*) \perp \text{Ran}(B^*)$, then $\|A + B\| = \max\{\|A\|, \|B\|\}$.

To obtain the result as we stated it, simply notice that $AB^* = 0$ implies $\text{Ran}(B^*) \subseteq \ker(A) = (\text{Ran}(A^*))^\perp$ so $\text{Ran}(B^*) \perp \text{Ran}(A^*)$. Similarly, $B^*A = 0$ implies that $\text{Ran}(A) \subseteq \ker(B^*) = (\text{Ran}(B))^{\perp}$ so $\text{Ran}(A) \perp \text{Ran}(B)$. Thus, if $AB^* = B^*A = 0$, we have $\text{Ran}(A) \perp \text{Ran}(B)$ and $\text{Ran}(A^*) \perp \text{Ran}(B^*)$ and one can then apply the theorem cited. **Q.E.D.**

Corollary 4.4.4. *The operators $\text{ff} + \text{bf}$, rec , $\text{bb} + \text{fb}$, and $\text{rr}_i^0 + \text{rr}_i^1 + \text{rc}_i$ (i fixed) are of norm 1.*

Proof. The proof is in fact reduced to a simple verification, by looking at the projections in the last algebra of the operators we defined, recalling that $\pi_a \pi_b = 0$ when $a \neq b$ and using the previous proposition.

Since $\pi_{\text{out}} \pi_{\text{in}} = \pi_{\text{in}} \pi_{\text{out}} = 0$, we have that $\text{bf} \times \text{ff}^*$ and $\text{ff}^* \times \text{bf}$ are equal to 0. Thus, using the preceding proposition, we have $\|\text{ff} + \text{bf}\| = \max\{\|\text{bf}\|, \|\text{ff}\|\} = 1$. A similar argument shows that $\|\text{fb} + \text{bb}\| = 1$.

Clearly, as $\pi_{ai} \pi_{bi} = \pi_{ao} \pi_{bo} = 0$ for $a \neq b$, we deduce by Proposition 4.4.3 that $\|\text{rec}\| = 1$.

Concerning $\text{rr}_i^0 + \text{rr}_i^1 + \text{rc}_i$, we fix $1 \leq i \leq p$ and notice that $\text{rc}_i \times (\text{rr}_i^0)^* = 0$ as a consequence of $\pi_{\text{mov-back}_i} \times \pi_{\text{back}_i} = 0$. Conversely, $(\text{rr}_i^0)^* \times \text{rc}_i = 0$ as a consequence of²⁷ $\pi_{\text{back}_{i+1}} \times \pi_{\text{mov-back}_i} = 0$. A similar argument shows that $\text{rc}_i \times (\text{rr}_i^1)^* = 0$ and $(\text{rr}_i^1)^* \times \text{rc}_i = 0$. Moreover, we know that $(\text{rr}_i^0)^* \times \text{rr}_i^1 = 0$ and $\text{rr}_i^1 \times (\text{rr}_i^0)^* = 0$ by just looking at the first term of the tensor in their definition. By an iterated use of Proposition 4.4.3, we get that $\|\text{rr}_i^0 + \text{rr}_i^1 + \text{rc}_i\| = \max\{\|\text{rr}_i^0\|, \|\text{rr}_i^1\|, \|\text{rc}_i\|\} = 1$. **Q.E.D.**

²⁶The range being the image or co-domain, i.e. $\text{Ran}(B) = \{x \mid \exists y B y = x\}$. The kernel is defined as $\ker(A) = \{x \mid A x = 0\}$.

²⁷When $i = p$, we consider that $\mathbf{q}_0 = \text{back}_{i+1}$ to ease the argument.

We are now in possession of all the material needed to study the 1-norm of the operator obtained by encoding of a deterministic pointer machine.

Proposition 4.4.5. *Let M be a DPM, s an initial configuration and $M_s^\bullet$ its encoding, then $\|M_s^\bullet\|_1 \leq 1$.*

Proof. Since we are working with deterministic machines, the transition relation is functional: for each $c \in C_M$ there is exactly one $t \in C_M$ —as the transition function is “total”— such that $c \rightarrow t$. Thus, Equation 4.10 becomes in this setting:

$$M_s^\bullet = \sum_{c \in C_M} \phi_{c,t} + \text{reject}_s$$

Since $M_s^\bullet$ is isomorphic to an element of $M_N(M_6(\mathfrak{R}))$, we will now compute $\|M_s^\bullet\|_1^N$. We first show that this matrix has at most one non-null coefficient in each column. Then we will use the fact that these coefficients are of norm at most 1.

To prove this we introduce the set $P = \{1, \dots, 6\}^p \times Q^\uparrow$ of *extended pseudo-configurations*, where $1, \dots, 6$ are a shorthand for $0i, \dots, \star o$. This set is a basis of the algebra of pseudo-states, and we write $M_s^\bullet = (a[c, c'])_{c, c' \in P}$. This way, the 1-norm of $M_c^\bullet$ is defined as

$$\|M_s^\bullet\|_1^N = \max_{c' \in P} \left\{ \sum_{c \in P} \|a[c, c']\| \right\}$$

Let $c = (a_1, \dots, a_n, \mathbf{q})$ be a (non-extended) pseudo-configuration. If c is a pseudo-configuration of the machine (i.e. $\mathbf{q}$ is an element of Q), then there is at most one pseudo-configuration t such that $c \rightarrow t$. This *atomic transition* can be:

Accept The operator $\phi_{c,t}$ is equal to 0 and the column is empty.

Reject The column corresponding to c contains only the operator $(\mathbf{q} \rightarrow \mathbf{back}_1)$ that encodes the transition $c \rightarrow (a_1, \dots, a_n, \mathbf{back}_1)$. And the norm of this operator is equal to 1.

Move forward the j th pointer and change state to $\mathbf{q}'$ The only extended pseudo-configurations introduced by the encoding is $\tilde{c} = (a_1, \dots, a_n, \mathbf{mov}_j^{\mathbf{q}, \mathbf{q}'})$. The column corresponding to c (resp. $\tilde{c}$) contains the operator $\text{ff} + \text{bf}$ (resp. rec), which is of norm 1 by Corollary 4.4.4.

Move backwards... This case is similar to the previous one.

Now we consider the extended pseudo-configurations, they are all introduced by the operator reject_s , and for $i = 1, \dots, p$ they are all of one of the following form:

$$\bar{c}_i^m = (a_1, \dots, a_n, \mathbf{mov-back}_i) \qquad \bar{c}_i^b = (a_1, \dots, a_n, \mathbf{back}_i)$$

The column corresponding to c (resp. $\bar{c}_i^b, \bar{c}_i^m$) contains only the operator encoding the transition from c to $\bar{c}_1^b$ (resp. the operator rm_i , the operator $\text{rr}_i + \text{rc}_i$), which is a norm 1 operator by Corollary 4.4.4.

We just showed that each column of the matrix contains at most one operator different from 0, and that this operator is always of norm 1. Hence, for any extended pseudo-configuration $c \in P$:

$$\sum_c a[c, c'] \leq 1$$

As a consequence of the definition of the 1-norm (Definition 3.4.6), we get

$$\|M_s^\bullet\|_1^N = \max_{c' \in P} \left\{ \sum_{c \in P} a[c, c'] \right\} \leq 1$$

Q.E.D.

Proposition 4.4.6.

$$\mathbf{L} \subseteq \{P_{+,1}\}$$

Proof. By Corollary 4.3.2 we know that $\mathbf{L} \subseteq \mathbf{DPM}$ and the previous proposition shows that the encoding of a DPM is of 1-norm inferior to 1, i.e. belongs to $P_{+,1}$. Q.E.D.

Theorem 4.4.7.

$$\mathbf{co-NL} = \{P_+\} = \{P_{\geq 0}\} \text{ and } \mathbf{L} = \{P_{+,1}\}$$

Proof. By combining Proposition 4.4.2, Proposition 4.4.6 and Corollary 3.5.4 from previous chapter. Q.E.D.

Results and perspectives

What have we done?

We presented several “pointers machines” to break the myth that “Pointers are log-space”: without any further assumption on the nature of the pointers and the way they are manipulated, this statement is not correct. JAG, PURPLE, SMM and KUM illustrate the way pointers may be handled as pebbles or data structure, without characterizing $\mathbf{L}$.

We then focused on a special kind of Finite Automata, that is considered as the better way to capture “pointer computation” as a Turing Machine would do it. We did not considered its alternating variant, but proved several classical claims about them. This was done to ease the introduction of Pointer Machines, which are similar on many points, but should rather be seen as universally non-deterministic. This machine model is introduced as a suitable model to be embedded in the operators.

The last section describes in its finest details this simulation, and encodes the basic operations of a PM so that they fulfil a constrain on their norm. The encoding of the rejection is somehow the more complicated part, for the only way to mimic rejection is to make the computation loop. Due to some oddity of the operator presentation, the simulation starts in parallel in all the possible initial states. This made the encoding of rejection even more pertinent: instead of a plain loop, the computation is started again from scratch, taking this time the actual initial pseudo-configuration.

What could we do?

This work is still very young and could be simplified in numerous ways.²⁸ We made some choices regarding the presentation, but the cursor between mathematics and computational complexity could be elsewhere. For instance, we could have defined the rejection of the PM as much closer to the

²⁸A first attempt using “unification algebra” can be found in a joint work with Marc BAGNOL [9].

way it is encoded in operators, that would have saved us the introduction of the “extended pseudo-configurations”, which is not very handy. Yet, we made this choice for it could look very surprising to define PM this way without having some insight on the way we are going to simulate their computation.

Regarding the future of this framework, we can mention —among numerous other, like the development of a higher-order programming language— the following ideas:

1. We gave the definition of 1-way FA not only to expose our culture, but because it should be clear that they can be encoded in operators all the same, with even a simpler presentation, for we would only need to define ff and frec to encode the movement of a head. But if we recall [Theorem 4.2.1](#), there are separation results both for non-deterministic and deterministic FA, 1-way and 2-ways. We know how to simulate every variation of FA: could we provide a new proof of those separation results? How could we reformulate the $\mathbf{L} = \mathbf{NL}$ conjecture? Is there any way to prove in a new setting $\mathbf{co-NL} = \mathbf{NL}$? In short, what is in our mathematical framework the meaning of those theorems and conjectures?
2. All the same, if we get rid of the tensor product and allow only one copy of $M_6(\mathbb{C})$, we describe the computation of a (one-head) Finite Automaton. The languages recognized by Finite Automata ($2\mathbf{DFA}(1)$ in our notation) corresponds to regular languages, which admits many other definitions, among them “being expressed using a regular expression”, or taking tools from the CHOMSKY hierarchy. But on the contrary, could a regular expression captures the computation of our observations?
3. What does the mathematical properties of operators teach us from a computational point of view? It is not very likely that $M_i^\bullet$ for M a 1-way FA has some mathematical properties that the encoding of a 2-way FA has not. All the same, there is a lot of isometries that go through our object, identifying a lot of them. Does that always corresponds to “reasonable” transformations between PM, or is it possible to determine a discrepancy between what is acceptable from a mathematical point of view, and what is acceptable from a computational point of view?
4. We took the dimension of the algebra of the pseudo-states to be bounded, but nothing restrain us from taking an infinite-dimensional algebra of state. What could that mean? Well, one way to grasp it would be to allow an infinite number of states in a PM. To what complexity class could correspond such a setting? It is natural to look in the direction of non-uniform complexity classes, and to ask ourselves what could be the difficulty to enumerate all the states of such a machine.

Conclusion and Perspectives

IN some sense, this work lacks imagination: we spent a lot of time carefully checking that all the components of our constructions were sound and well defined. The understanding of those objects requires some work, and we have the feeling that the work of simplification we led for proof circuits could be pushed further, and applied to operators. We hope this work will help to get a better understanding of those wonderful objects.

Proof circuits do not seem to offer that much perspectives: it is a practical way to define properly the parallel computation of Proof Nets, as it is perfectly suited to mimic Boolean circuits. But it seems at the same time that this framework is too restrictive: there is no clear way to handle types other than boolean, and it is really pertinent to study them only for constant-depth resources. It could nevertheless be interesting to build proof circuits that actually solves some problems: maybe the type constraints could tell us something about the necessity of increasing the depth of the formulae?

On the other hand, the opportunities to push the study of operators forward seem endless: for instance, how could we build a concrete algorithm as operator? The framework may not look very handy, but in fact the properties of the projections does all the hard work. This could open new perspectives, by bounding more precisely the resources: JONES [84] showed using a pointer-like programming language that a constant factor was a significant in the case of *actual* computation, when we come back from theoretical models to computation with computers. This work was pushed further by BEN-AMRAM and JONES [17] and proposes another way of handling the complexity, by paying more attention to what is hidden inside this $\mathcal{O}$ -notation. We believe our operators could offer a framework to handle such concerns.

The study of the logical part of this work could also be developed: what could be the proof of an observation? Could we adapt this construction to other logics? The tensor logic of MELLIÈS and TABAREAU [104] could allow us to bridge the gap with categorical semantics, in a framework where we already know precisely how to handle resources. This logic moreover allows to express sequentiality without being forced to renounce to proof nets, among other tools provided by Linear Logic. This could also offer a nice alternative to λ -calculi to study Boolean proof nets from a sequential perspective. One could also benefit from the work led by SEILLER [121], to approach GoI with combinatorial tools.

There are other perspectives we could mention, but we prefer to conclude with an intriguing remark: we saw that graphs were heavily linked to log-space computation. On the other hand, matrix are usually related to Boolean circuits, for they are the ideal input to process in parallel (you may think of the multiplication of two matrices, a problem that is in $\mathbf{NC}^1$). And yet we modelised non-deterministic log-space computation with matrices, and parallel computation with graphs. Could we try to apply operators to proof circuits, and vice-versa?



Operator Algebras Preliminaries

OPERATOR algebra can be seen as an “harmonious blend of algebra and analysis” (CURIEN [30]). It is also a complex subject, where spectral theory, functional analysis, linear algebra and topology meet. Each one of those perspectives came with its own notation, interpretation and nomenclature.¹

So we felt it could be useful to synthesize in this first appendix the material needed to fully understand the construction of Chapter 3 and the encoding of Section 4.4. We tried to use a clear and neutral notation —summed up in the table of notations, p. xiii.

Operators are taken here to be continuous linear functions on a separable Hilbert space, and operators algebras are taken to be self-adjoint. It makes us focus on C^* -algebras and von Neumann algebras, and we will go straightforward to the definitions of those notions, in order to rapidly define projections, factors, hyperfiniteness and type. Those notions are the one in use in our construction, and even if it is not mandatory to be an expert on this subject to understand it, we felt it necessary to ease their introduction by some reminders.

The first section defines von Neumann algebras following two paths: the first one will set up all the material needed to understand them as an algebraic object whereas the second offers a topological approach. The two definitions we propose —Definition A.1.18 and Definition A.1.22— are of course completely equivalent, thanks to the famous “double commutant theorem”. We left aside the so-called “abstract” definition of von Neumann algebras, *a.k.a.* W^* -algebras, as C^* -algebras that are the dual space of a Banach algebra.

The second section exposes some basis regarding von Neumann algebras, explain their links with matrices, and develop the quite complex operation of crossed product. This last part involves massively groups and is useful to understand the normative pair we use in our construction (cf. Corollary 3.4.4).

Almost everything in this appendix is called a definition, but in fact it mixes definitions, lemmas and theorems. This appendix is in debt to some notes written by CURIEN [30] and by GIRARD [53], which prepared and eased the lecture of more complex handbooks. Thomas SEILLER patiently explained me

¹Even the definition of operators is subject to variations! A textbook on Operator Theory and Operator Algebra [109, p. 3] defines operators as “the bounded linear maps from a normed vector space to itself”, a Linear Algebra handbook [12, p. 57] defines them as “linear map[s] from a vector space to itself” whereas a handbook of Theory of Operator Algebra [86, p. 2] takes them to be mappings between two vector spaces. Here, we will constraint operators to be bounded linear maps from a normed vector space to itself.

a lot regarding those notions, and the reader may find a quick overview of the needed material in the appendix of one of his article [120]. His thesis [121] presents extensively this subject and constitutes an excellent bridge toward those objects. Some motivation and a different way of exposing the subject can also be found in GIRARD [50, Appendix on C*-algebra, pp. 48–60]. All those references could help the reader more used to Linear Logic to adapt to operator algebra.

The reader who would like to go further could refer to the textbooks of CONWAY [26] for the bases of the theory (starting from Hilbert and Banach spaces, C*-algebra, topologies on them), of MURPHY [109] for an excellent introduction to the theory of operator algebras. The series of TAKESAKI [126, 127, 128] remains the most complete reference.

A.1 Towards von Neumann Algebras

There is not much to say about prerequisites needed to read this section, for we tried to define everything in a simple manner. Just agree that we use δ for the *Kronecker delta*, defined by

$$\delta_{ij} = \begin{cases} 0, & \text{if } i \neq j \\ 1, & \text{if } i = j \end{cases}$$

The algebraic way

The reader used to vector spaces, operators and C*-algebra, may directly go to Definition A.1.18 of von Neumann algebras. Elsewhere, (s)he can have a look at Figure A.1 to see the summary of the algebraic structures we need to properly defines von Neumann algebras.

Definition A.1.1 (Distance, Metric space). A *metric space* is a pair (S, d) with $S \neq \emptyset$ a set and $d : S^2 \mapsto \mathbb{R}$ a *distance over S* such that $\forall x, y, z \in S$

$$\begin{aligned} d(x, y) &= d(y, x) && \text{(Symmetry)} \\ d(x, y) &= 0 \text{ iff } x = y && \text{(Separation)} \\ d(x, z) &\leq d(x, y) + d(y, z) && \text{(Subadditivity)} \end{aligned}$$

Definition A.1.2 (Algebraic Structures). Let E be a set, $+$ and $\times$ two binary operations, we define the following set of rules:

$$\begin{aligned} \forall a, b \in E, (a + b) &\in E && \text{(Closure)} \\ \forall a, b, c \in E, (a + b) + c &= a + (b + c) && \text{(Associativity)} \\ \exists e_+ \in E, \forall a \in E, e_+ + a &= a + e_+ = a && \text{(Identity element)} \\ \forall a \in E, \exists -a \in E, a + (-a) &= (-a) + a = e_+ && \text{(Inverse element)} \\ \forall a, b, c \in E, a + b &= b + a && \text{(Commutativity)} \\ \forall a_1, a_2, a_3 \in E, (a_1 + a_2) \times a_3 &= (a_1 \times a_3) + (a_2 \times a_3) && (\times \text{ is right-distributive over } +) \\ \forall a_1, a_2, a_3 \in E, a_1 \times (a_2 + a_3) &= (a_1 \times a_2) + (a_1 \times a_3) && (\times \text{ is left-distributive over } +) \end{aligned}$$

The rules are instantiated to $+$, but a simple substitution makes clear what are the rules for $\times$.

This set of rules allow us to define several key-concepts at the same time:

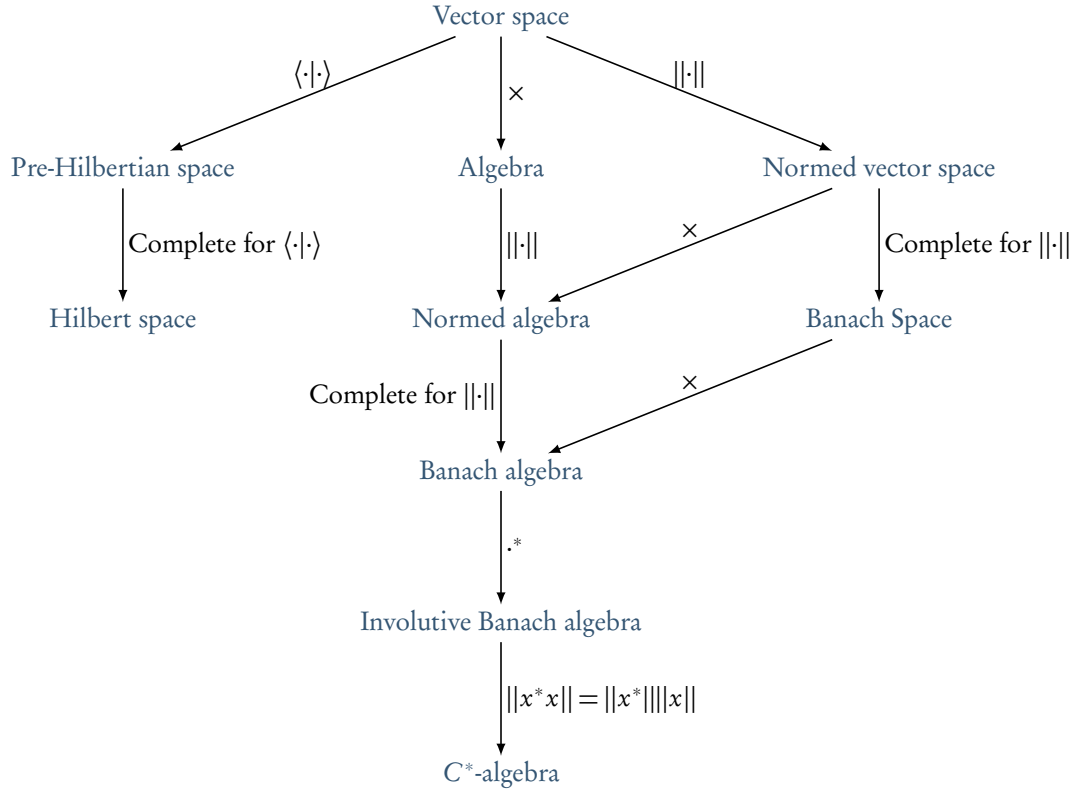


Figure A.1: Dependency of the algebraic structures

- If $\mathcal{G} = (E, +)$ respects eqs. (Closure) to (Inverse element), then $\mathcal{G}$ is a *group*. If $\mathcal{G}$ moreover respects eq. (Commutativity), then $\mathcal{G}$ is said to be *commutative* (or *abelian*).
- If $R = (E, +, \times)$ is such that $(E, +)$ is a commutative group, $\times$ respects eq. (Closure), eq. (Identity element) and is distributive over $+$,² then R is a *pseudo-ring*. If moreover $\times$ is associative, R is a *ring* and if moreover $\times$ is commutative then R is a *commutative ring*.
- If $\mathcal{F} = (E, +, \times)$ is such that $(E, +)$ and $(E \setminus \{e_+\}, \times)$ are two commutative groups and $\times$ distributes over $+$, then $\mathcal{F}$ is a *field* (also called a *corpus*).

In all of those case, we call E the *underlying set*. For the sake of simplicity, we will denote the underlying set of a group $\mathcal{G}$ (resp. a field $\mathcal{F}$) with $\underline{\mathcal{G}}$ (resp. $\underline{\mathcal{F}}$). We will also use later on this notation for vector spaces.

The operations $+$ and $\times$ are mostly referred to as *addition* and *multiplication*, and the neutral elements e_+ and $e_\times$ are most of the time denoted 0 and 1 . The inverse of $a \in E$ under $+$ is usually written $-a$ whereas the notation a^{-1} is kept for the inverse under $\times$.

Of particular interest to us will be the group of permutations over a set.

²Meaning that $\times$ is left-distributive over $+$ and $\times$ is right-distributive over $+$.

Definition A.I.3 (Permutation of a set, Symmetric group). Let S be a set, a *permutation* of S is a bijection from S to itself. We set $\mathcal{S}(S)$ to be the group of all permutations of S , we write $\mathcal{S}_n$ if $S = \{1, \dots, n\}$ and $\mathcal{S}$ if $S = \mathbb{N}$.

For $<$ a total ordering over S and for all $\sigma \in \mathcal{S}(S)$, we define its *signature* $\text{sgn}(\sigma)$ to be the parity of the number of inversions for σ

$$\text{sgn}(\sigma) := \text{Card}(\{(x, y) \mid x < y \text{ and } \sigma(x) > \sigma(y)\}) \pmod{2}$$

Definition A.I.4 (Subgroup). Let $\mathcal{G} = (E, +)$ be a group, $\mathcal{H} = (F, +|_F)$ is a *subgroup* of $\mathcal{G}$ if $F \subseteq E$ and $+|_F$ is a group operation on F . We write $\mathcal{H} \leq \mathcal{G}$. If $F \subsetneq E$, $\mathcal{H}$ is said to be a *proper subgroup* of $\mathcal{G}$. If $\forall h \in F, \forall x \in E, x + h + x^{-1} \in F$, $\mathcal{H}$ is a *normal subgroup* of $\mathcal{G}$ and we write $\mathcal{H} \triangleleft \mathcal{G}$.

We introduce in the following two operations, the vector addition and the vector multiplication, and denote also them with $+$ and $\times$. This is quite classical and harmless, since the type and the context should always make clear which one is in use. Moreover, the algebras axioms tells us basically that those operations are the same over different sets.

Definition A.I.5 (Vector space, Algebra over a field). A *vector space* $\mathbf{V}$ over a field $\mathcal{F} = (\mathcal{F}, +, \times)$ is a set $\underline{\mathbf{V}}$ endowed with two binary operations, $+: \underline{\mathbf{V}}^2 \mapsto \underline{\mathbf{V}}$ the vector addition and $\cdot: \mathcal{F} \times \underline{\mathbf{V}} \mapsto \underline{\mathbf{V}}$ the scalar multiplication, such that $(\underline{\mathbf{V}}, +)$ is a commutative group and such that for all *vectors* $u, v, w \in \underline{\mathbf{V}}$, for all *scalars* $\lambda, \mu \in \mathcal{F}$, the following holds:

$$\begin{aligned} \lambda \cdot (u + v) &= (\lambda \cdot u) + (\lambda \cdot v) && (\cdot \text{ is left-distributive over } +) \\ (\lambda + \mu) \cdot u &= (\lambda \cdot u) + (\mu \cdot u) && (\cdot \text{ is right-distributive over } +) \\ (\lambda \times \mu) \cdot u &= \lambda \cdot (\mu \cdot u) && (\text{Distributivity of } \times \text{ over } \cdot) \\ e_{\times} \cdot u &= u && (\text{Identity element of } \times) \end{aligned}$$

We say that $\mathbf{W}$ is a *linear subspace* of $\mathbf{V}$ over $\mathcal{F}$ if $\underline{\mathbf{W}} \subseteq \underline{\mathbf{V}}$, $\underline{\mathbf{W}} \neq \emptyset$ and $\forall x, y \in \underline{\mathbf{W}}, \forall \lambda \in \mathcal{F}, x + y \in \underline{\mathbf{W}}$ and $\lambda x \in \underline{\mathbf{W}}$.

If moreover $\mathbf{V}$ is endowed with a binary operation $\times$ such that

$$\begin{aligned} (u + v) \times w &= (u \times v) + (u \times w) && (\times \text{ is right-distributive over } +) \\ u \times (v + w) &= (u \times v) + (u \times w) && (\times \text{ is right-distributive over } +) \\ (\lambda \cdot u) \times (\mu \cdot v) &= (\lambda \times \mu) \cdot (u \times v) && (\text{Compatibility with scalars}) \end{aligned}$$

Then $\mathcal{A} = (\underline{\mathbf{V}}, +, \cdot, \times)$ is an *algebra over the field* $\mathcal{F}$. If moreover $\times$ is associative, i.e. if $u \times (v \times w) = (u \times v) \times w$, then $\mathcal{A}$ is said to be an *associative algebra* over $\mathcal{F}$.

We say that $\mathcal{U}$ is a *subalgebra* of $\mathcal{A}$ over $\mathcal{F}$ if $\underline{\mathcal{U}} \subseteq \underline{\mathcal{A}}$, $\mathcal{U}$ is a linear subspace of $\mathcal{A}$ over $\mathcal{F}$ and $\forall u, v \in \underline{\mathcal{U}}, u \times v \in \underline{\mathcal{U}}$.

Definition A.I.6 (Basis of a vector space). A basis $\mathbb{B}$ of a $\mathcal{F}$ -vector space $\mathbf{V}$ is a linearly independent subset of $\underline{\mathbf{V}}$ that spans $\mathbf{V}$.

$\mathbb{B} \subseteq \underline{\mathbf{V}}$ is *linearly independent* if for every $\mathbb{B}_0 = \{f_1, \dots, f_n\} \subset \mathbb{B}$, $\forall (\lambda_1, \dots, \lambda_n) \in \mathcal{F}^n$, $(\lambda_1 \cdot f_1) + \dots + (\lambda_n \cdot f_n) = e_+$ implies $\lambda_1 = \dots = \lambda_n = e_+$.

$\mathbb{B} = (f_i)_{i \in I}$ spans $\mathbb{V}$ if $\forall u \in \mathbb{V}$ there exists $(\lambda_i)_{i \in I} \in \mathcal{F}^I$ with $\{i \in I \mid \lambda_i \neq 0\}$ finite such that $u = \sum_{i \in I} \lambda_i f_i$.

All bases of a vector space $\mathbb{V}$ have the same cardinality, called the *dimension* of $\mathbb{V}$ and denoted by $\dim(\mathbb{V})$.

Definition A.1.7 (Inner product, orthogonality, pre-Hilbertian space). Let $\mathbb{V}$ be a $\mathbb{C}$ -vector space, an *inner product* is a map $\langle \cdot | \cdot \rangle : \mathbb{V}^2 \mapsto \mathbb{C}$, such that $\forall x, y, z \in \mathbb{V}, \lambda \in \mathbb{C}$

$$\begin{aligned} \langle x | y \rangle &= \overline{\langle y | x \rangle} && \text{(Conjugate symmetry)} \\ \langle \lambda x | y \rangle &= \lambda \langle x | y \rangle && \text{(Linearity in the first argument)} \\ \langle x + y | z \rangle &= \langle x | z \rangle + \langle y | z \rangle \\ \langle x | x \rangle &\geq 0 \text{ and } \langle x | x \rangle = 0 \text{ implies } x = 0 && \text{(Positive-definiteness)} \end{aligned}$$

Where $\overline{(\cdot)}$ is the conjugate defined by $\overline{a + ib} = a - ib$. We call this product *sesquilinear*.³

If $\langle x | y \rangle = 0$, we say that x and y are orthogonal and we define $W^\perp$ as the *orthogonal complement* of a subspace W of $\mathbb{V}$ by $W^\perp = \{x \in \mathbb{V} \mid \forall y \in W, \langle x | y \rangle = 0\}$.

A *pre-hilbertian space* $(E, \langle \cdot | \cdot \rangle)$ is a vector space E with an inner product $\langle \cdot | \cdot \rangle$.

Definition A.1.8 (Norm, normed vector space). Let $\mathbb{V}$ be a $\mathcal{F}$ -vector space, a *norm over $\mathbb{V}$* is a map $\|\cdot\| : \mathbb{V} \mapsto \mathbb{R}$ such that $\forall u, v \in \mathbb{V}, \forall \lambda \in \mathcal{F}$, the following holds:

$$\begin{aligned} \|u\| &= 0_{\mathbb{R}} \text{ iff } u = e_+ && \text{(Separation)} \\ \|u + v\| &\leq \|u\| + \|v\| && \text{(Subadditivity)} \\ \|\lambda \cdot u\| &= |\lambda| \cdot \|u\| && \text{(Positive homogeneity)} \end{aligned}$$

A *normed vector space* is a pair $(\mathbb{V}, \|\cdot\|)$ where $\mathbb{V}$ is a vector space and $\|\cdot\|$ is a norm on $\mathbb{V}$. The norm induces a distance by letting $d(x, y) = \|x - y\|$.

Definition A.1.9 (Convergent sequence, Cauchy sequence). A sequence $x_1, x_2, \dots$ of a metric space (S, d) converge (in S) if $\exists l \in S$ —called *the limit*— such that $\forall \epsilon \in \mathbb{R}^{>0}, \exists N \in \mathbb{N}, \forall n \in \mathbb{N}$,

$$n \geq N \Rightarrow d(x_n - l) < \epsilon$$

We write $x_1, x_2, \dots \rightarrow l$. We say that $x_1, x_2, \dots$ is a *Cauchy sequence* if $\forall \epsilon \in \mathbb{R}^{>0}, \exists n \in \mathbb{N}, \forall p > n, \forall q \in \mathbb{N}$,

$$d(x_{p+q} - x_p) < \epsilon$$

Every convergent sequence is a Cauchy sequence, and (S, d) is called *complete* (or a *Cauchy space*) if every Cauchy sequence in S converges.

This definition can easily be adapted to normed vector space: we only have to replace the distance by the norm.

Definition A.1.10 (Pointwise convergent). Let $\{f_n\}$ be a sequence of functions with the same domain and co-domain, $\{f_n\}$ converges *pointwise* to f iff for all x in the domain, $\lim_{n \rightarrow \infty} f_n(x) = f(x)$.

³We take the mathematical convention, our inner product is linear in the first argument and anti-linear in its second. The physicists use the opposite convention.

Definition A.1.11 (Banach space). A *Banach space* is a normed vector space which is complete for the induced distance.

Definition A.1.12 (Normed algebra, Banach algebra). A *normed algebra* $(S, +, \cdot, \times, \|\cdot\|)$ is

$(S, +, \cdot, \times)$ an associative algebra over $\mathbb{R}$ or $\mathbb{C}$,

$(S, \|\cdot\|)$ a normed vector space

such that the algebra multiplication and the norm respects: $\forall u, v \in S, \|u \times v\| \leq \|u\| \times \|v\|$.

In particular, if $(S, \|\cdot\|)$ is a Banach space, $(S, +, \cdot, \times, \|\cdot\|)$ is a *Banach algebra*. If $\times$ admit an identity element $e_\times$, and $\|e_\times\| = 1_{\mathbb{R}}$, $(S, +, \cdot, \times, \|\cdot\|)$ is *unital*.

Definition A.1.13 (Involutive Banach algebra, C^* -algebra). If a Banach algebra $\mathcal{B} = (S, +, \cdot, \times, \|\cdot\|)$ admits a map $(\cdot)^* : S \rightarrow S$ such that $\forall u, v \in S, \lambda \in \mathbb{C}$:

$$(\lambda \cdot u)^* = \bar{\lambda} \cdot u^* \quad (u + v)^* = u^* + v^* \quad (u \times v)^* = v^* \times u^* \quad u^{**} = u \quad \|u^*\| = \|u\|$$

then V is called an *involutive Banach algebra* and the map $(\cdot)^*$ is called the *involution*. If the involution satisfies the following additional condition:

$$\|u^* u\| = \|u^*\| \|u\| = \|u\|^2$$

then $\mathcal{B}$ is called a C^* -algebra (or *stellar algebra*).

A **-homomorphism* between two C^* -algebras $\mathcal{B}$ and $\mathcal{C}$ is a bounded linear map $f : \mathcal{B} \rightarrow \mathcal{C}$ that “preserve $(\cdot)^*$ ”, i.e. $\forall x, y \in \mathcal{B}, f(x^*) = f(x)^*$ and $f(xy) = f(x)f(y)$.

If f is injective, it is an isometry, and if f is bijective, it is a C^* -isomorphism.

Definition A.1.14 (Hilbert space). A *Hilbert space* $\mathbb{H}$ is a $\mathbb{C}$ -vector space with a norm $\langle \cdot | \cdot \rangle$ such that $\mathbb{H}$ is a complete metric space with respect to the distance d defined by

$$d(x, y) = \|x - y\| = \sqrt{\langle x - y | x - y \rangle}$$

It is a Banach space where $\forall x \in \mathbb{H}, \langle x | x \rangle = \|x\|^2$.

We always take our Hilbert spaces to be infinite-dimensional and separable. That induces the existence of an orthonormal separable basis, but the usual concept of basis (defined in Definition A.1.6) cannot be applied “as it” to infinite Hilbert spaces. In this setting, we use *Hilbert basis* that allow us to approximate every element with an infinite sequence of finite elements.

Definition A.1.15 (Hilbert basis). A family $\mathbb{F} = (x_i)_{i \in I} \subseteq \mathbb{H}$ is a *Hilbert basis* of $\mathbb{H}$ if

$$\begin{aligned} \forall i, j \in I^2, \langle x_i | x_j \rangle &= \delta_{ij} & (\mathbb{F} \text{ is orthonormal}) \\ \forall x \in \mathbb{H}, \exists \{\lambda_i\}_{i \in I} \in \mathbb{C}^I, x &= \sum_{i \in I} \lambda_i x_i & (\mathbb{F} \text{ is complete}) \end{aligned}$$

We will speak of “homomorphism” between two algebraic structure as long as this function preserve the algebraic structure under treatment. For instance, given $(\mathcal{G}, +)$ and $(\mathcal{G}', +')$ two groups, $f : \mathcal{G} \rightarrow \mathcal{G}'$ is an homomorphism if $\forall g_1, g_2 \in \mathcal{G}, f(g_1 + g_2) = f(g_1) +' f(g_2)$.

Among the homomorphisms, of particular interest to us are the homomorphisms between vector spaces, called *linear maps*.

Definition A.I.16 (Linear map). Let V and W be two $\mathcal{F}$ -vector spaces, f a map from V to W is said to be a *linear map* (or is $\mathcal{F}$ -linear) iff

$$\forall x, y \in \underline{\mathcal{F}}, \forall \mu \in \underline{V}, f(x + (\mu \cdot y)) = f(x) + (\mu \cdot f(y))$$

Let $\mathcal{L}(V, W)$ be the set of linear applications from V to W , we define the *dual space* V^* to be $\mathcal{L}(V, \mathcal{F})$.

If $V = W$, f is said to be an *endomorphism* and we write $\mathcal{L}(V)$ the set of linear maps from V to V . For all V a $\mathcal{F}$ -vector spaces, $\mathcal{L}(V)$ endowed with the pointwise addition and the composition of functions as a multiplication operation is an algebra over the field $\mathcal{F}$.

In fact, we instantiate the previous definition to the special case of the Hilbert space, and restrict ourselves to the case we will use to define *operators* as the bounded mappings between two Hilbert spaces.

Definition A.I.17 (Operator, $\mathcal{B}(\mathbb{H})$). A linear map u from $\mathbb{H}$ to itself is *bounded* if there exists $c \in \mathbb{N}$ such that $\forall x \in \mathbb{H}, \|u(x)\|_{\mathbb{H}} \leq c\|x\|_{\mathbb{H}}$. We define $\mathcal{B}(\mathbb{H})$ to be the set of bounded linear maps from $\mathbb{H}$ to itself, and we define a norm on $\mathcal{B}(\mathbb{H})$:

$$\|u\|_{\mathcal{B}(\mathbb{H})} := \sup_{x \in \mathbb{H}} \left(\frac{\|ux\|_{\mathbb{H}}}{\|x\|_{\mathbb{H}}} \right)$$

For all $u \in \mathcal{B}(\mathbb{H})$, there exists a unique $u^* \in \mathcal{B}(\mathbb{H})$, the *adjoint of u* , such that $(\cdot)^*$ is an involution satisfying $\|u^*u\| = \|u\|^2$.

If we take the addition to be the addition of vectorial spaces, the multiplication to be the composition and $\cdot$ to be the multiplication by a scalar, then $(\mathcal{B}(\mathbb{H}), +, \times, \cdot, \|\cdot\|, (\cdot)^*)$ is a C^* -algebra.

It is equivalent for an operator to be bounded or to be continuous, see [Definition A.I.20](#).

Definition A.I.18 (Commutant, von Neumann algebra). Let $\mathfrak{M} \subset \mathcal{B}(\mathbb{H})$ be a C^* -subalgebra, we define the *commutant* of $\mathfrak{M}$ (relatively to $\mathcal{B}(\mathbb{H})$), but we will omit this precision) to be the set

$$\mathfrak{M}' := \{u \in \mathcal{B}(\mathbb{H}) \mid \forall m \in \mathfrak{M}, mu = um\}$$

If $\mathfrak{M} = \mathfrak{M}'$, $\mathfrak{M}$ is a von Neumann algebra.

The topological way

This path will be way shorter, mainly because we only need weak operator topology in the following. Some of the notions previously defined (as norm, Hilbert space, or convergence) are required for this section. We will spend some time with a classical example midway between topology and algebra.

Definition A.I.19 (Topology, Topological space). Let X be a set and $\mathcal{T} \subseteq \mathcal{P}_{\text{fin}}(X)$ a family of subset of X . We say that $\mathcal{T}$ is a *topology* on X if:

$$X \in \mathcal{T}, \emptyset \in \mathcal{T} \quad \forall x_1, x_2 \in \mathcal{T}, x_1 \cap x_2 \in \mathcal{T} \quad \forall I, \forall (x_i)_{i \in I} \in \mathcal{T}^I, \cup_{i \in I} x_i \in \mathcal{T}$$

If $\mathcal{T}$ is a topology on X then the pair $(X, \mathcal{T})$ is a *topological space*. The elements of $\mathcal{T}$ are called *open sets* in X . The subset $Y \subseteq X$ is said to be *closed* if its complement is in $\mathcal{T}$.

For all $x \in X$, the *neighbourhood of x* is a set $N \subseteq X$ such that there exist an open set O with $x \in O \subseteq N$. We write $N(x)$ the set of the neighbourhoods of x . We call X *compact* if for all $\{U_\alpha\}_{\alpha \in A}$ of open subsets of X such that $X = \bigcup_{\alpha \in A} U_\alpha$, there is a finite $J \subseteq A$ such that $X = \bigcup_{i \in J} U_i$. We say that X is *locally compact*, if $\forall x \in X, N(x)$ is compact.

Definition A.1.20 (Continuous map). Let $(X, \mathcal{T})$ and $(X', \mathcal{T}')$ be two topological spaces, a map $f : X \rightarrow X'$ is *continuous* if and only if for every $x \in X$ and every $V' \in \mathcal{N}(f(x))$ there exists $V \in \mathcal{N}(x)$ such that $y \in V$ implies $f(y) \in V'$.

Said differently, f is continuous if for every open set $O \subseteq X'$, $f^{-1}(O) = \{x \in X \mid f(x) \in O\}$ is an open subset of X .

Definition A.1.21 (Weak topology on $\mathbb{H}$ and $\mathcal{B}(\mathbb{H})$). Let $\mathbb{H}$ be a Hilbert space, we say that a sequence $\{x_i\}_{i \in \mathbb{N}}$ converges weakly to 0 when $\langle x_i | y \rangle \rightarrow 0$ for all $y \in \mathbb{H}$. Weak convergence is thus a point-wise or direction-wise convergence.

On $\mathcal{B}(\mathbb{H})$, we says that $\{u_i\}_{i \in \mathbb{N}}$ converges weakly to 0 when, for any $x \in \mathbb{H}$, $u_i x$ converges weakly to 0. This defines the *weak operator topology*.

We can show that $\mathcal{B}(\mathbb{H})$ is the dual of a space denoted by $\mathcal{B}(\mathbb{H})_*$ containing the *trace-class operators*. For further details, the reader may refer to MURPHY [109, section 2.4, pp. 53 sqq.] or TAKESAKI [126, p. 63]. We recall this result only to define the σ -weak topology: if A is a topological space and A^* is its dual, the *weak** topology on A^* is defined as the point-wise topology.

Definition A.1.22 (von Neumann algebra (2)). A von Neumann algebra $\mathfrak{M}$ is a C^* -subalgebra of $\mathcal{B}(\mathbb{H})$ which is closed in the weak operator topology and contains the identity operator.

We are beginning to mix topology and linear algebra with the following example of one of the most famous Hilbert space. The following construction could be generalized with the Lebesgue spaces L^2 over measurable spaces, i.e. the set of measurable functions whose power of 2 has finite integral.

Example A.1.1 ($\ell^2(S)$). Let S be a set, we define $\ell^2(S)$ to be the set of square-summable functions

$$\ell^2(S) = \{f : S \rightarrow \mathbb{C} \mid \sum_{s \in S} |f(s)|^2 < \infty\}$$

where $|\cdot|$ is the absolute value of the complex number, i.e. if $f(s) = a + ib$, $|f(s)| = \sqrt{a^2 + b^2}$.

This set $\ell^2(S)$ is a $\mathbb{C}$ -vectorial space if we take for the addition the pointwise addition of the functions, for the multiplication the multiplication of $\mathbb{C}$. It is a Hilbert space if we define an inner product $\forall f, g \in \ell^2(S)$

$$\langle f | g \rangle = \sum_{s \in S} f(s) \overline{g(s)}$$

where $\overline{g(s)} = \overline{g(s)}$ is the conjugate of $g(s)$. The norm is as usual defined as $\|f\| = \sqrt{\langle f | f \rangle}$ and we can easily check that $\ell^2(S)$ is complete with respect to that norm.

This example is interesting because every Hilbert space with an Hilbert basis $\mathbb{F}$ is isomorphic to $\ell^2(\mathbb{F})$. In particular, if $\mathbb{F}$ is a discrete group $\mathcal{G}$, $\ell^2(\mathcal{G})$ is a von Neumann algebra.

And if $\mathcal{G}$ is the set bijections from $\mathbb{N}$ to $\mathbb{N}$ that permute a finite number of elements, the von Neumann algebra obtained is a II_1 -factor, whose definition follows.

A.2 Some Elements on von Neumann Algebras

In this section we go one step further by introducing some more complex notions. We recall some of the terminology of the von Neumann algebras, explains the links between operator algebra and matrices algebra, and then study the crossed-product of a von Neumann algebra with a group.

Basics of von Neumann algebras

First, we should remark that von Neumann algebras also admit a *tensor product*, which comes basically from the tensor product for Hilbert spaces.

Definition A.2.1 (Tensor product for the von Neumann algebras). If for $i = 1, 2$, $\mathfrak{M}_i$ is a von Neumann algebra on the Hilbert space $\mathbb{H}_i$, then $\mathfrak{M}_1 \otimes \mathfrak{M}_2$ is the von Neumann algebra on $\mathbb{H}_1 \otimes \mathbb{H}_2$.

Much of the terminology of C^* -algebras is the same as for the operator theory on Hilbert space.

Remark (Operators zoo). Let $u \in \mathcal{B}(\mathbb{H})$, we say that u is

Normal if $uu^* = u^*u$

Self-adjoint or Hermitian if $u^* = u$

A projection if $u = u^* = u^2$

A partial isometry if $uu^* = p$ for p a projection

If $\mathcal{B}(\mathbb{H})$ is *unital* (it admits an identity operator $\text{Id}_{\mathcal{B}(\mathbb{H})}$), we say that u is

Unitary if $uu^* = u^*u = \text{Id}_{\mathcal{B}(\mathbb{H})}$

An isometry if $uu^* = \text{Id}_{\mathcal{B}(\mathbb{H})}$

Definition A.2.2 (Center of a von Neumann algebra, factor). The *center* of a von Neumann algebra $\mathfrak{M}$ is defined as $\{x \in \mathfrak{M} \mid \forall y \in \mathfrak{M}, xy = yx\}$. This is the same as the set $\mathfrak{M} \cap \mathfrak{M}'$ for $\mathfrak{M}'$ the commutant of $\mathfrak{M}$.

A von Neumann algebra $\mathfrak{N}$ is a *factor* if its center is trivial, i.e. if it consists only of multiples (scalar operators) of the identity operator.

Every von Neumann algebra on a separable Hilbert space is isomorphic to a direct integral (a continuous direct sum) of factors, and this decomposition is essentially unique.

So the study of von Neumann algebras is partially subsumed by the study of classes of factors. The factors themselves can be classified by their sets of projections, whose definition follows. Informally, a projection sends the Hilbert space $\mathbb{H}$ onto a closed subspace of $\mathbb{H}$. Depending on the dimensions and the properties of this closed subspace, we may define an equivalence relation on them.

Definition A.2.3 (Finite and infinite projections). A projection is an operators $p \in \mathcal{B}(\mathbb{H})$ such that $p = p^* = p^2$. We write $\Pi(\mathfrak{M})$ the set of projections in $\mathfrak{M}$. Let p, q be two projections of a von Neumann algebra, if $qp = pq = p$ or equivalently if $\dim(p) \subseteq \dim(q)$, then we write $p \leq q$.

In a von Neumann algebra, we can define an equivalence relation on $\Pi(\mathfrak{M})$ by $p \sim_{\mathfrak{M}} q$ if there exists a partial isometry $u \in \mathfrak{M}$ such that $uu^* = p$ and $u^*u = q$. This equivalence relation is sometimes called the *Murray and von Neumann equivalence*.

If there exists q such that $q < p$ and $q \sim p$, p is said to be *infinite*. Conversely, a projection p is *finite* if $\forall q$ a projection, $q \sim p$ and $q < p$ implies $q = p$. A projection p is said to be *minimal* if $q \leq p$ implies $q = 0$ or $q = p$.

A von Neumann algebra can be “recreated” from its projections, it is sufficient to know $\Pi(\mathfrak{M})$ to describe totally $\mathfrak{M}$. In fact, there is a bijection between the projections of $\mathfrak{M}$ and the subspaces of $\mathbb{H}$ that belong to it. Projections also allow to classify them, whenever they have minimal projections or not, finite projections or not. We will only use the II_1 factor, so we define only this one, and we define at the same time hyperfiniteness.

Definition A.2.4 (Hyperfiniteness type II_1 factor). A von Neumann algebra $\mathfrak{N}$ where $\leq$ is total is a *factor*, and it is said to be of type II if $\mathfrak{N}$ contains finite projections but has no minimal projections. If the identity of $\mathfrak{N}$ is a finite projection, $\mathfrak{N}$ is of type II_1 .

A von Neumann algebra $\mathfrak{M}$ is *hyperfiniteness* if there exists an ascending sequence $(\mathfrak{M}_i)_{i \in \mathbb{N}}$ of finite-dimensional subalgebras of $\mathfrak{M}$ such that their union $\cup_{i \in \mathbb{N}} \mathfrak{M}_i$ is dense in $\mathfrak{M}$.⁴ Equivalently,

$$\dim(\mathfrak{M}_i) < \infty \text{ for all } i \quad \mathfrak{M}_1 \subset \mathfrak{M}_2 \subset \mathfrak{M}_3 \dots \quad \mathfrak{M} = \cup_{i \in \mathbb{N}} \mathfrak{M}_i$$

The type II_1 hyperfiniteness factor is unique up to isomorphism and we will write it $\mathfrak{R}$.

The hyperfiniteness of a factor $\mathfrak{M}$ implies that the operators of $\mathfrak{M}$ can be approximated⁴ by matrices, i.e. by operators acting on a finite-dimensional Hilbert space.

On matrices

Matrices shares a lot with hyperfiniteness von Neumann algebras, partly thanks to the possibility they have to represent linear functions. We will also see how to represent matrices algebra with operator algebra. But first, let's agree on some notations.

Definition A.2.5 (Matrices). Let $M_{n,m}(\mathcal{F})$ be the set of matrices with coefficients in $\mathcal{F}$, with n rows and m columns. If $n = m$, the matrices are said to be *square* and we write $M_n(\mathcal{F})$ the set of square matrices. Most of the time we won't specify the field where the matrix takes its coefficients, and we define as usual:

- I_n to be the *identity matrix* of size n , 0_n to be the *null matrix*.
- $M[i, j]$ to be the entry in the i th row and j th column of M .
- If $\forall i, j \leq n, M[i, j] = 0$ if $i \neq j$, M_n is a *diagonal matrix*.
- If $\forall i, j \leq n, M[i, j] = M[j, i]$, M_n is a *symmetric*.
- The *trace* of a matrix M_n is defined by $\text{tr}(M_n) = \sum_{i=1}^n M[i, i]$.
- A matrix M_n is a *projection* if $M_n^2 = M_n$.
- The *determinant* of a matrix M_n is defined as usual by the Leibniz formula:

$$\det(M_n) = \sum_{\sigma \in \mathcal{S}_n} \text{sgn}(\sigma) \prod_{i=1}^n m_{i, \sigma(i)}$$

- If $M_1[i, j] = M_2[j, i]$, we say that M_2 is the *transpose* of M_1 and we write $M_1^t = M_2$.

⁴For the σ -weak topology that we did not define.

- If $\overline{M[i, j]} = M[i, j]$, we say that $\overline{M}$ is the *conjugate* of the matrix M .
- If $M^* = \overline{M}^t$, we say that M^* is the *conjugate-transpose* of M .
- If M_n is such that $\exists d \in \mathbb{N}^*$ and $M_n^d = 0_n$, M_n is said to be *nilpotent* and the smallest such d is the *degree (of nilpotency) of M_n* .

We should remark that $\forall n \in \mathbb{N}$, the set of matrices $M_n(\mathbb{C})$ is a non-commutative algebra over $\mathbb{C}$. We simply takes the classical matrix addition and multiplication, the scalar multiplication, and $M_n(\mathbb{C})$ inherits all the properties of $\mathbb{C}$ as an algebra.

Definition A.2.6 (Matrix representation of linear maps). Let V be a $\mathcal{F}$ -vector space of dimension n and let $\mathbb{B} = \{b_1, \dots, b_n\}$ be a basis of V and $u \in \mathcal{L}(V)$.

$$\forall 1 \leq k \leq n, \exists \lambda_1^k, \dots, \lambda_n^k \in \mathcal{F} \text{ such that } u(b_k) = \lambda_1^k(b_1) + \dots + \lambda_n^k(b_n)$$

And we define

$$[u]_{\mathbb{B}} := \begin{pmatrix} \lambda_1^1 & \dots & \lambda_1^n \\ \vdots & \ddots & \vdots \\ \lambda_n^1 & \dots & \lambda_n^n \end{pmatrix}$$

Remark that

- u is nilpotent ($u \circ \dots \circ u = 0$) iff $[u]_{\mathbb{B}}$ is nilpotent.
- If u is nilpotent, $\det([u]_{\mathbb{B}}) = \text{tr}([u]_{\mathbb{B}}) = 0$.
- For all $\mathbb{B}_1$ and $\mathbb{B}_2$ two basis of V , $\text{tr}([u]_{\mathbb{B}_1}) = \text{tr}([u]_{\mathbb{B}_2})$. So we can speak of *the* trace of a linear map.

In the case of the II_1 hyperfinite factor, any element can be approximated by finite matrices, and we can always define a “normalised trace” by $\text{Tr}(M_n) = \frac{\text{tr}(M_n)}{n}$. This trace allows us to define a generalization of the determinant, the *Fuglede-Kadison determinant*.

But for now we consider the reverse operation, i.e. representing any matrix algebra in the II_1 hyperfinite factor.

Definition A.2.7 (Embedding matrices algebra in the II_1 hyperfinite factor). Let $\mathfrak{R}$ be the II_1 hyperfinite factor and $\Pi(\mathfrak{R})$ be the set of projections of $\mathfrak{R}$. As $\mathfrak{R}$ is of type II, it has no minimal projection, so we can always decompose a projection, i.e. for all $p \in \Pi(\mathfrak{R})$, there exists $q_1, q_2 \in \Pi(\mathfrak{R})$ such that $q_1 q_2 = 0$ and $q_1 + q_2 = p$. Moreover, we can chose q_1 and q_2 such that $q_1 \sim q_2$ i.e. such that they act on subspaces with the same dimensions.

If we generalize, $\forall p \in \Pi(\mathfrak{R})$, for all $n \in \mathbb{N}^*$, there exists $q_1, \dots, q_n \in \Pi(\mathfrak{R})$ such that the following holds:

$$q_i \sim q_j \text{ for all } 1 \leq i \leq j \leq n \quad \sum_{i=1}^n q_i = p \quad q_i q_j = \delta_{ij} q_i$$

Now we take $p_1, \dots, p_n \in \Pi(\mathfrak{R})$ such that $\sum_{i=1}^n p_i = 1$. As $p_i \sim p_j$, we know that there exists $u_{i,j}$ a partial isometry such that

$$u_{i,j} u_{i,j}^* = p_j \quad u_{i,j}^* u_{i,j} = p_i$$

Such a partial isometry is not unique, and we may chose the $u_{i,i+1}$ such that

$$u_{i,j} = \begin{cases} u_{i,i+1} \cdots u_{j-1,j} & \text{if } i < j \\ u_{j,j+1}^* \cdots u_{i-1,i}^* & \text{if } i > j \\ p_i & \text{if } i = j \end{cases}$$

And we have that

$$u_{i+1,i} = u_{i,i+1}^* \quad (\text{A.1})$$

$$u_{i,j} u_{k,m} = \delta_{jk} u_{i,m} \quad (\text{A.2})$$

Now we define the embedding Θ from the set of $n \times n$ matrices over $\mathbb{C}$ to the Π_1 hyperfinite factor $\mathfrak{R}$. Let M be a $n \times n$ matrix, we define $\Theta(M[i, j]) = M[i, j] u_{i,j}$, i.e. $\Theta(M) = \sum_{i,j}^n M[i, j] u_{i,j}$.

Let M_n and M'_n be two matrices over $\mathbb{C}$, we check that Θ commutes with the operations of the algebra. There is not much to check, because Θ basically inherits of all the algebraic properties of $\mathbb{C}$, so we just illustrate with $+$, $\times$ and the transposition.

$$\begin{aligned} (\Theta M_n) + (\Theta M'_n) &= \left(\sum_{i,j}^n M[i, j] u_{i,j} \right) + \left(\sum_{i,j}^n M'[i, j] u_{i,j} \right) \\ &= \sum_{i,j}^n ((M[i, j] u_{i,j}) + (M'[i, j] u_{i,j})) \\ &= \sum_{i,j}^n (M[i, j] + M'[i, j]) u_{i,j} \\ &= \Theta(M_n + M'_n) \end{aligned}$$

$$\begin{aligned} (\Theta M_n) \times (\Theta M'_n) &= \left(\sum_{i,j}^n M[i, j] u_{i,j} \right) \times \left(\sum_{i',j'}^n M'[i', j'] u_{i',j'} \right) \\ &= \sum_{i,j}^n \sum_{i',j'}^n (M[i, j] u_{i,j} \times M'[i', j'] u_{i',j'}) \\ &= \sum_{i,j}^n \sum_k^n (M[i, k] \times M'[k, j]) u_{i,j} \quad (\text{By Equation A.2}) \\ &= \Theta \left(\sum_{i,j}^n \sum_k^n (M[i, k] \times M'[k, j]) \right) \\ &= \Theta(M_n \times M'_n) \end{aligned}$$

$$\begin{aligned} \Theta(M_n^t) &= \sum_{i,j}^k (M[j, i] u_{i,j}) \\ &= \sum_{i,j}^k (M[j, i] u_{j,i}^*) \quad (\text{By Equation A.1}) \\ &= \Theta(M_n)^* \end{aligned}$$

So we just checked that Θ is an isomorphism, and we can check that the properties of the normalised trace is preserved by this isomorphism.

In fact Θ depends on the size of n , so we should write $\Theta_{n+1} : M_n(\mathbb{C}) \rightarrow \mathfrak{K}$, and we should also remark that such an embedding is not unique.

Groups and von Neumann algebras

We will in the following show how to internalize the action of a group into a von Neumann algebra. We need first some preliminaries definitions.

Definition A.2.8 (Topological group). Let $\mathcal{G} = (E, +)$ be a group endowed with a topology $\mathcal{T}$, $\mathcal{G}$ is a *topological group* if $+$ and $(\cdot)^{-1}$ are continuous maps.

A topological group is a *compact group* if it is compact, any compact group is locally compact.

Definition A.2.9 (Action of a group). Let $\mathcal{G} = (E, +)$ be a group and X be a set, an action of $\mathcal{G}$ on X is a function $\alpha : E \times X \rightarrow X$ such that for every $a, b \in E$ and $x \in X$, $\alpha(e_+, x) = x$ and $\alpha(a, \alpha(b, x)) = \alpha(a + b, x)$.

Equivalently, if we set $\mathcal{S}(X)$ to be the group of all permutations of X , an action α of $\mathcal{G}$ on X is a homomorphism $\alpha : E \rightarrow \mathcal{S}(X)$.

An action of a topological group $\mathcal{G}$ on a von Neumann algebra $\mathfrak{M}$ is a continuous homomorphism of $\mathcal{G}$ into $\text{Aut}(\mathfrak{M})$, the set of automorphisms of $\mathfrak{M}$.

Definition A.2.10 (Representations). Let $\mathfrak{M}$ be a von Neumann algebra. A couple $(\mathbb{H}, \rho)$ where $\mathbb{H}$ is a Hilbert space and ρ is a $*$ -homomorphism from $\mathfrak{M}$ to $\mathcal{B}(\mathbb{H})$ is a *representation* of $\mathfrak{M}$. If ρ is injective, we say the representation is *faithful*.

HAAGERUP [68] proved that among the numerous representations of a von Neumann algebra, there was a *standard representation*, a representation satisfying several important properties.

The operation that will be of interest to us will be that of taking the *crossed product* of an algebra and a group. This operation is closely related to that of semi-direct product of groups and is a way of internalizing automorphisms. Given an algebra $\mathfrak{M}$ and a group $\mathcal{G}$ of automorphisms of $\mathfrak{M}$, we construct the algebra $\mathfrak{M} \rtimes \mathcal{G}$ generated by the elements of $\mathfrak{M}$ and the elements of $\mathcal{G}$.

Definition A.2.11 (Crossed product (representations)). Let $(\mathbb{H}, \rho)$ be a representation of a von Neumann algebra $\mathfrak{M}$, $\mathcal{G}$ a locally compact group, and α an action of $\mathcal{G}$ on $\mathfrak{M}$. Let $\mathbb{K} = L^2(\mathcal{G}, \mathbb{H})$ be the Hilbert space of square-summable $\mathbb{H}$ -valued functions on $\mathcal{G}$. We define representations π_α of $\mathfrak{M}$ and λ of $\mathcal{G}$ on $\mathbb{K}$ as follows:

$$\begin{aligned} (\pi_\alpha(u).x)(g) &= (\rho(\alpha(g)^{-1}(u))x)(g) & (u \in \mathfrak{M}, x \in \mathbb{H}, g \in \mathcal{G}) \\ (\lambda(g).x)(h) &= x(g^{-1}h) & (g, h \in \mathcal{G}, x \in \mathbb{K}) \end{aligned}$$

Then the von Neumann algebra on $\mathbb{K}$ generated by $\pi_\alpha(\mathfrak{M})$ and $\lambda(\mathcal{G})$ is called the crossed product of $(\mathbb{H}, \rho)$ by α .

An important fact is that the result of the crossed product does not depend on the chosen representation of $\mathfrak{M}$.

Theorem A.2.1 (TAKESAKI [127, Theorem 1.7, p. 241]). Let $(\mathbb{H}, \rho)$ and $(\mathbb{K}, \rho')$ be two faithful representations of a von Neumann algebra $\mathfrak{M}$, and let $\mathcal{G}$ be a locally compact group together with an action α on $\mathfrak{M}$. Then there exists an isomorphism between the crossed product of $(\mathbb{H}, \rho)$ by α and the crossed product $(\mathbb{K}, \rho')$ by α .

As a consequence, one can define *the* crossed product of a von Neumann algebra and a group acting on it by choosing a particular representation. Of course, the natural choice is to consider the *standard representation*.

Definition A.2.12 (Crossed product). Let $\mathfrak{M}$ be a von Neumann algebra, $\mathcal{G}$ a group and α an action of $\mathcal{G}$ on $\mathfrak{M}$. The algebra $\mathfrak{M} \rtimes_{\alpha} \mathcal{G}$ is defined as the crossed product of the standard representation of $\mathfrak{M}$ by α .

A particular case of crossed product is the crossed product of $\mathbb{C}$ by a (trivial) action of a group $\mathcal{G}$. The resulting algebra is usually called the group von Neumann algebra $\mathfrak{N}(\mathcal{G})$ of $\mathcal{G}$. As it turns out, the operation of internalizing automorphisms of algebras (the crossed product) and the operation of internalizing automorphisms of groups (the semi-direct product) correspond: the algebra $\mathfrak{N}(\mathcal{G} \rtimes_{\alpha} H)$ is isomorphic to $\mathfrak{N}(\mathcal{G}) \rtimes_{\tilde{\alpha}} H$ where $\tilde{\alpha}$ is the action of H on $\mathfrak{N}(\mathcal{G})$ induced by the action of H on $\mathcal{G}$.

An Alternate Proof of $\text{co-NL} \subseteq \text{NPM}$

Contents

A.1	Towards von Neumann Algebras	126
	The algebraic way	
	The topological way	
A.2	Some Elements on von Neumann Algebras	132
	Basics of von Neumann algebras	
	On matrices	
	Groups and von Neumann algebras	

THIS short appendix is here to recollect the first version [10] of the proof that $\text{co-NL} \subseteq \text{NPM}$, a part of Corollary 4.3.2 according to the notation in this thesis. It is quite different to the one given in Chapter 4 for it does not use the characterization of NL provided by FA, but exhibits a NPM that solves a co-NL -complete problem and a *transducer* that perform reductions between any co-NL problem and the problem we solve. We develop a “pointer reduction” a bit complex and yet interesting, for it was not previously developed as far as we know.

This proof was written when we were not aware of the existence of FA: it is in some sense more direct, but it does not present our PMs as a member of the large family of “reasonable” computational model. We will use the most common method¹: first we define and solves a co-NL -complete problem with a NPM, and then define a “pointer reduction”. Of course a co-NL problem is also a NL problem, but once again we won’t take this equality for granted.

Numerous elements of Chapter 1 and Chapter 4 are needed to understand this appendix, for we won’t recall for instance the definitions of reductions (Definition 1.1.8 and Definition 1.1.9), complete problem (Definition 1.1.10), computation graph (Definition 1.1.5), NPM (Definition 4.3.1). Some of the claims stated in Section 4.2 will be used “as it”, even if they were proved for FA and not for PM, we consider the proofs could be adapted really easily.

¹That can be reminded to the reader for instance in ARORA and BARAK [4, pp. 88–89].

B.1 Solving a co-NL-Complete Problem

We are going to solve the most common **co-NL**-complete problem, the complementary of *stconn* (Problem B). It is reasonable to modify *stconn* so that the path we are looking for is always from the node numbered 1 to the “last” node, i.e. the one with the highest number, which corresponds to the size of the graph.

We define the set

$$STConnComp = \{k \in \mathbb{N} \mid k \text{ encode a directed graph of size } n \text{ with no path from } 1 \text{ to } n\}$$

As we will produce the actual transition function that solves this problem, we should be precise regarding the encoding of a graph we adopt: given a graph of size n , the input will be

$$\underbrace{\star 00 \dots 00}_{n \text{ bits}} \boxed{1} \overbrace{a_{11} 0 a_{12} 0 \dots 0 a_{1n-1} 0 a_{1n}}^{\text{edges going from 1}} \boxed{1} \dots \boxed{1} \overbrace{a_{n1} 0 a_{n2} 0 \dots 0 a_{nn-1} 0 a_{nn}}^{\text{edges going from } n} \boxed{1}$$

where (a_{ij}) is the adjacency list, i.e. $a_{ij} = 1$ if there is an edge from the node numbered by i to the node numbered by j , 0 elsewhere. The boxed bits in the figure above are “separating” bits, between the unary encoding of n and the adjacency list, and between the coding of the edges going from i to the coding of the edges going from $i + 1$. This will allow us a trick in our algorithm: when we read the bit following a_{ij} , if it is a 0 then the next bit corresponds to a_{ij+1} , elsewhere it corresponds to a_{i+11} .

Proposition B.1.1. *There exists $M \in \text{NPM}(5)$ that decides $STConnComp$.*

Proof. We define M such that $M_s(k)$ with $s = \{\star, \star, \star, \star, \text{Init}\}$ accepts iff $k \in STConnComp$.

The transition relation of M is presented in Figure B.1. Informally, our algorithm goes as follow:

- We have a counter on the size of the path followed, p_1 . It moves right on the unary expression of n every time an edge is found (on the second line of Equation B.10). If it reads a 1 after moving (Equation B.12), it means we already followed n edges, so we accept because if there is a path from 1 to n then there exists a path of size at most n . Note that p_1 does not move in any other transition, except for the initialisation.
- The couple p_2, p_3 acts as follow: when p_2 is reading a_{ij} , p_3 is at a_{j1} . If $a_{ij} = 1$ (premise of Equation B.10), a non-deterministic transition takes place: on one way we continue to scan the outgoing edges from i , on the other we increment p_1 , place p_2 at a_{j1} and p_3 at a_{11} .
- The last pointer p_4 is here to check when $a_{ij} = 1$ if $j = n$, it also stays on the unary encoding of n . If p_4 reaches a 1 when p_2 reads that there is an edge, it means that there is an edge whose target is n , and so we reject (Equation B.11).
- When p_2 finishes to browse the adjacency list of an edge, we accept (Equation B.6).

We can make easily sure that $M_s(k)$ accepts iff $k \in STConnComp$, elsewhere $M_s(k)$ rejects, i.e. $M_s(k)$ always halts. As defined, M is not a NPM, for it does not fulfil all the requisites we listed in Section 4.3. But we can easily “complete” the transition function: for all $c \in C_M$ such that $c \rightarrow \emptyset$, add $c \rightarrow \text{accept}$. We may also reasonably modify it so that at most one head move at each transition. Q.E.D.

$$\begin{aligned}
(\star, \star, \star, \star, \text{Init}) &\rightarrow (p_1+, p_2+, p_3+, p_4+, \text{Init}) & (\text{B.1}) \\
(\star, 0, \star, \star, \text{Init}) &\rightarrow (\epsilon_1, p_2+, p_3+, \epsilon_4, \text{Init}) & (\text{B.2}) \\
(\star, 1, \star, \star, \text{Init}) &\rightarrow (\epsilon_1, p_2+, \epsilon_3, \epsilon_4, \text{out.edge?}) & (\text{B.3}) \\
(\star, 0, \star, \star, \text{out.edge?}) &\rightarrow (\epsilon_1, p_2+, \epsilon_3, p_4+, \text{no.edge}) & (\text{B.4}) \\
(\star, 0, \star, \star, \text{no.edge}) &\rightarrow (\epsilon_1, \epsilon_2, p_3+, \epsilon_4, \text{p3.next.node}) & (\text{B.5}) \\
(\star, 1, \star, \star, \text{no.edge}) &\rightarrow \text{accept} & (\text{B.6}) \\
(\star, \star, \star, \star, \text{p3.next.node}) &\rightarrow (\epsilon_1, \epsilon_2, p_3+, \epsilon_4, \text{reading.sep.bit}) & (\text{B.7}) \\
(\star, \star, 0, \star, \text{reading.sep.bit}) &\rightarrow (\epsilon_1, \epsilon_2, p_3+, \epsilon_4, \text{p3.next.node}) & (\text{B.8}) \\
(\star, \star, 1, \star, \text{reading.sep.bit}) &\rightarrow (\epsilon_1, p_2+, \epsilon_3, \epsilon_4, \text{out.edge?}) & (\text{B.9}) \\
(\star, 1, \star, \star, \text{out.edge?}) &\rightarrow \begin{cases} (\epsilon_1, p_2+, \epsilon_3, p_4+, \text{no.edge}) \\ (p_1+, \epsilon_2, \epsilon_3, p_4+, \text{edge.found}) \end{cases} & (\text{B.10}) \\
(\star, \star, \star, 1, \text{edge.found}) &\rightarrow \text{reject} & (\text{B.11}) \\
(1, \star, \star, 0, \text{edge.found}) &\rightarrow \text{accept} & (\text{B.12}) \\
(\star, \star, \star, 0, \text{edge.found}) &\rightarrow (\epsilon_1, p_2-, \epsilon_3, p_4-, \text{rewind.p2.p4}) & (\text{B.13}) \\
(\star, \star, \star, [\text{not } \star], \text{rewind.p2.p4}) &\rightarrow (\epsilon_1, p_2-, \epsilon_3, p_4-, \text{rewind.p2.p4}) & (\text{B.14}) \\
(\star, \star, \star, \star, \text{rewind.p2.p4}) &\rightarrow (\epsilon_1, p_2-, \epsilon_3, \epsilon_4, \text{rewind.p2}) & (\text{B.15}) \\
(\star, [\text{not } \star], \star, \star, \text{rewind.p2}) &\rightarrow (\epsilon_1, p_2-, \epsilon_3, \epsilon_4, \text{rewind.p2}) & (\text{B.16}) \\
(\star, \star, \star, \star, \text{rewind.p2}) &\rightarrow (\epsilon_1, p_2+, p_3-, \epsilon_4, \text{exchange.p2.p3.}) & (\text{B.17}) \\
(\star, \star, [\text{not } \star], \star, \text{exchange.p2.p3.}) &\rightarrow (\epsilon_1, p_2+, p_3-, \epsilon_4, \text{exchange.p2.p3.}) & (\text{B.18}) \\
(\star, \star, \star, \star, \text{exchange.p2.p3.}) &\rightarrow (\epsilon_1, \epsilon_2, p_3+, \epsilon_4, \text{get.p3.to.start}) & (\text{B.19}) \\
(\star, \star, 0, \star, \text{get.p3.to.start}) &\rightarrow (\epsilon_1, \epsilon_2, p_3+, \epsilon_4, \text{get.p3.to.start}) & (\text{B.20}) \\
(\star, \star, 1, \star, \text{get.p3.to.start}) &\rightarrow (\epsilon_1, p_2+, \epsilon_3, \epsilon_4, \text{out.edge?}) & (\text{B.21})
\end{aligned}$$

Recall that $\star$ (resp. $[\text{not } \star]$) is any symbol among $\{0, 1, \star\}$ (resp. $\{0, 1\}$).

Figure B.1: The transition relation that decides *STConnComp*

B.2 Pointer-Reduction

So now, we have to prove that we can reduce any **co-NL**-problem to *STConnComp* using only pointers. We will adapt the classical log-space implicitly computable function that reduces any problem in **NL** to *stconn*. Given a **co-NL**-problem P , there exists a universally non-deterministic log-space Turing Machine M that decides it. To solve P is just to establish if there is **no** path in the computational graph of M from the initial configuration to a rejecting configuration.

We will in the following:

- shows how PMs can express “any integer” and perform some computation,
- explains the properties of the encoding of $G_{M(n)}$ as an integer,
- prove that a DPM can “compute” $G_{M(n)}$ given M and n ,
- explain how to compose PMs, and that will complete the demonstration.

CLAIM 8: *A PM can express any power of the size of the input.*

Proof. Let fix an input n of size $|n|$, our claim is that given j pointers, a PM can count up to $|n|^j$. We use once again the additional pointers as the hands of a clock: we make the 1st pointer go right, and every time the i th pointer made a round-trip (that is, is back on $\star$), the $i + 1$ th pointer goes one cell right. Clearly, we can express any distance inferior to $|n|^j$. **Q.E.D.**

We will for the sake of simplicity consider that any distance inferior to $|n|^j$ can be expressed by a single pointer, even if it may require several pointer to be properly expressed.

CLAIM 9: *A PM can compute addition, subtraction, power, division, multiplication, modulo and binary logarithm.*

Proof. We can adopt the same proof as the one we developed for [Claim 5](#) regarding the FA. To compute $\log$, just remember that $\log(x) = y$ iff $x = 2^y$, so our PM can put a pointer (that will be an estimation for y) on a cell, computes its power of two, and looks if it is the same value as x . If it is, the value we are looking for is found, elsewhere it just make the pointer move one cell right and makes again this computation. **Q.E.D.**

Given any problem $P \in \text{co-NL}$, there exists a universally non-deterministic Turing Machine M such that M accepts n^2 iff $n \in P$. We assume that P is a decision problem, that M works on alphabet $\{0, 1\}$, has one read-only tape and one read-write working tape whose precise bound is $k \times (\log(|n|))$. We may also assume that the name of the states is written in binary, so that for $|Q| = q$ the number of states of M , any state may be written with $\lceil \log(q) \rceil$ bits. At last, we may assume that the instructions to move the heads are written with two bits. All those assumptions make clear that M may be entirely described with a binary string.

We remind the reader that the computation graph $G_{M(n)}$ is finite and acyclic and that its size is the number of configuration of $M(n)$. We can take M to always halt, so $M(n)$ rejects iff there is a branch that reaches **reject**, and we can also assume that only one configuration has for state **reject**: this will allow us to test only for the existence of a path in $G_{M(n)}$ from the initial configuration to that configuration.

²Meaning that **all** branches reach **accept** after a finite number of transitions.

We know that $M(n)$ has less than

$$2^{(k \times (\log(|n|)))} \times (k \times (\log(|n|))) \times (\log(|n|)) \times \lceil \log(q) \rceil$$

different configurations. It reflects respectively the content of the working tape, the position of the read-write and the read-only heads and the state. This is equivalent to $2^{O(\log(|n|))}$, so we know there exists a d such that $M(n)$ has less than $|n|^d$ different configurations.

Any configuration of $M(n)$ may be described as

$$\underbrace{0100011110 \dots 010110}_{\text{Position of the read head}} \underbrace{\sigma 0 \sigma 0 \dots \sigma 0 \sigma 1 \sigma 0 \dots \sigma 0 \sigma 0 \sigma 0 \sigma 0 \sigma 0 \sigma 0}_{\text{Working tape and position of the working head}} \underbrace{001 \dots 10}_{\text{state}}$$

where the $\lceil \log(|n|) \rceil$ first bits encode in binary the position of the reading head, the symbols σ correspond to the bits on the working tape, the bit that follows σ equals 1 iff the working head is on that cell. The remaining $\lceil \log(q) \rceil$ bits express the current state.

This binary string is of length $\lceil \log(|n|) \rceil \times (2 \times (k \times \lceil \log(|n|) \rceil) \times \lceil \log(q) \rceil)$, i.e. there exists a e such that this string is of length inferior to $e \times \log(|n|)^2$.

Among all the binary string of size $e \times \log(|n|)^2$, some of them corresponds to configurations, and some of them don't (for instance because the working head is supposed to be in several places at the same time), we will call them "phantom configurations".

Lemma B.2.1 (Pointer-reduction). *For all universally non-deterministic log-space Turing Machine M , there exists a DPM T such that for all n , given a pointer p_d with $\#p_d = j$, T accepts if the j th bit of the encoding of $G_{M(n)}$ is 1, rejects if it is 0.*

Proof. Recall we use the encoding exposed earlier to express the encoding of $G_{M(n)}$ the computation graph of M on input n . Let T proceed as follows:

- It computes the number of binary strings of size $e \times \log(|n|)^2$. This number is bounded by $2^{(e \times \log(|n|))^2}$ and we saw previously that a PM could express such distances. Then T compares this value to j : if j is inferior, it rejects, if j is equal, it accepts, elsewhere it goes on. This reflects the n initial bits set to 0 to express in unary the size of the graph.
- Elsewhere T establishes if j corresponds to a "separating bit" and accepts or rejects accordingly, that can be simply made with the division and modulo operations.
- Elsewhere, j encodes a query regarding the presence or absence of transition between two configurations a and b . If $a = b$, there is no need to explore this transition,³ and T rejects. Elsewhere T establishes if there is a transition between a and b , and accepts or rejects accordingly.

This last point needs to be made more precise: if $j > 2^{e \times \log(|n|)^2}$ and if j does not corresponds to a "separating bit", it means that the value of j corresponds to the absence or presence of an edge between two nodes. So there exists a and b such that $j = a_{ab}$. A simple arithmetic of pointers allows us to retrieve those two values expressed as integers (i.e. as a distance).

Then, they are converted to binary strings: the positions of the read-only heads needs a bit of pointer arithmetic to be obtained and compared, but the rest of the integer just needs to be compared

³Because that would imply that there is a transition from a configuration to itself, and so $M(n)$ is stuck in a loop.

bitwise. The rest of the binary expression of the node encodes directly the configuration, and as all the transitions makes only local changes in them, there is only a constant number of information to be remembered.

Every time there is a difference between the binary expression of a and the binary expression of b , T checks that the difference between them is legal regarding the transition function of M —that may be encoded in the states of T or may be given as a parameter.

The transducer T also have to check that a and b are not “phantom configurations” and that j is not “too big”, i.e. does not represent a query on nodes that does not exists. **Q.E.D.**

One should notice that even if this computation is quite violent, for the number of pointer of T will be *very* high, we did not used non-determinacy, so T is a DPM.

Corollary B.2.2.

$$\text{co-NL} \subseteq \text{NPM}$$

Proof. Let $P \in \text{co-NL}$, there exists a universally non-deterministic Turing Machines N such that $\mathcal{L}(N) = P$. Suppose given $n \in \mathbb{N}$, we will compose the NPM M that solves $STConnComp$ with the transducer T that compute $G_{M(n)}$.

Every time M has to read a value, it asks T by letting a pointer be on the position j of the value it wants to know. There is some kind of layer of abstraction in this composition, for M go through the input tape without ever reading the actual values, but asks the values to T , which actually read n .

We have to make sure that the j of Lemma B.2.1 can be big enough: what is the size of the encoding of $G_{N(n)}$? We encode the graph as being of size $2^{e \times \log(|n|)^2}$, i.e. we also take “phantom configurations” to be nodes of $G_{N(n)}$. The encoding of this “completed” graph—for every string of size $e \times \log(|n|)^2$ is taken to be one of its node, even if it is not reachable—is of size $\mathcal{O}((2^{\log(|n|)^2})^2)$, an expression bounded by a power of $|n|$, so we can express it.

We can suppose moreover that there is a transition between the “phantom configuration” encoded by 000...001 and the initial configuration, and that there exists a transition between any rejecting configuration and the “phantom configuration” encoded by 111...111. This allows us to keep the $STConnComp$ algorithm as it is, computing only if there is no path from the node 1 to the node n .

The transducer T can compute the configuration graph of $N(x)$ bit-by-bit and pass it to N which solves $STConnComp$. So $M \circ T(n)$ accepts iff there is no path from 1 to an rejecting configuration in $G_{N(n)}$, i.e. iff $N(n)$ accepts. So $P \in \text{NPM}$. **Q.E.D.**

Index

A

Algebra of states, 82, III
Alphabet, 9, 97

B

Banach Algebra, 59, 127
Binary Integer, 68
Boolean
 circuit, 22
 function, 22
 Proof net, 30

C

Complexity classes, xvi, 9
 $\text{STA}(*, *, *)$, 9
 AC^i , II, 24
 DFA, 99, 103
 L, 10, 120
 NC^i , II, 24
 NFA, 99, 103
 NL, 10
 PCC^i , 40
 P, 10
 bPCC^i , iii, 40
 co-NL , 10, 118
 mBN^i , 40
 $\{P_{+,1}\}$, 120
 $\{P_{+}\}$, 118
 $\{P_{\geq 0}\}$, 118
Computation path, 7
Configuration
 Bounds on the number of configurations,
 8, 102, 145
 Pseudo-configuration, 108, III

Successor configuration, 7

Contraction, 62
Crossed product, 81, 137
Curry-Howard correspondence, 3, 14, 60
Cut, 26
Cut-elimination, 26, 40

D

Direct Connection Language, 49
 for Boolean circuits, 23
 for Boolean proof nets, 32

E

Exponentials, 60

F

Finite Automata, 98

G

Graph
 Computation graph, 8, 144
 Families of computation graphs, 46
 Encoding of a graph, 142
 Thick graph, 65
Group, 127
 Group action, 60
 Permutation group, 82, 128

H

Head
 Read-only head, 6, 98, 144
 Read-write head, 105
 Sensing head, 101

I

Invariance thesis, 4
 Parallel case, 5
 Isometry, 60

J

JAG, 94, 95

K

KUM, 94, 96

L

Lambda-calcul, 59
 Lambda-term, 14, 63, 68
 Linear Algebra, 15, 16
 Linear Logic, 12
 Bounded Linear Logic, 18, 24
 Elementary Linear Logic, 18, 58, 68
 Intuitionist Multiplicative Linear Logic,
 50
 Multiplicative Linear Logic, 13, 26, 60
 Unbounded Multiplicative Linear Logic,
 24, 26
 log-space, 11, 48, 97

M

Matrix, 60, 62, 65, 70, 83, 134

N

Nilpotency, 58, 80–83, 89
 Norm, 127, 129, 130
 1-norm, 82, 87–89, 119
 Normative pair, 81, 82

O

Observation, 82
 Operator algebra, 60
 Operators, 58–60

P

Parity, 32
 Permutations, 54, 60

Piece (of a proof circuit), 33, 42

Pointer, 50, 94, 96

Pointer machine, 90, 94, 106

Polynomial time, 11

Problem, 9

 Complete problem, 10

see also Source-Target Connectivity, Undirected

 Source-Target Connectivity

Projection, 60, 66, 67, 74, 78, 111, 113, 133, 135

Proof circuit, 35

Proof net, 14, 25, 50, 61, 63

Proof-as-program, *see* Curry-Howard
 correspondence

Pseudo-states, 110

PURPLE, 94, 95

R

Reduction, 10

 log-space reduction, 10

 between problems, 10

 Pointer-reduction, 145

S

Simulation, 4

SMM, 94, 96

Source-Target Connectivity, 12

 Complementary, 142

T

Tensor product, 133

Transposition, 16

U

Undirected Source-Target Connectivity, 12, 22,
 29, 43, 44

V

Vector Space, 127, 128

von Neumann algebra, 58–60, 81, 83, 112, 131

W

Word, 6, 9, 23, 96

Bibliography

- [1] Eric W. ALLENDER, Michael C. LOUI, and Kenneth W. REGAN. “Complexity Classes”. In: *Algorithms and Theory of Computation Handbook*. Edited by Mikhail J. Atallah and Marina Blanton. CRC, 1998, pages 674–696. DOI: [10.1201/9781584888239](https://doi.org/10.1201/9781584888239) (cited on p. 4).
- [2] Eric W. ALLENDER, Michael C. LOUI, and Kenneth W. REGAN. “Reducibility and Completeness”. In: *Algorithms and Theory of Computation Handbook*. Edited by Mikhail J. Atallah and Marina Blanton. CRC, 1998, pages 697–724. DOI: [10.1201/9781584888239](https://doi.org/10.1201/9781584888239) (cited on p. 4).
- [3] Jean-Marc ANDREOLI, Roberto MAIELI, and Paul RUET. “Non-commutative proof construction: A constraint-based approach”. In: *Annals of Pure and Applied Logic* 142.1–3 (2006), pages 212–244. DOI: [10.1016/j.apal.2006.01.002](https://doi.org/10.1016/j.apal.2006.01.002) (cited on p. 54).
- [4] Sanjeev ARORA and Boaz BARAK. *Computational Complexity: A Modern Approach*. Volume 1. Cambridge University Press, 2009. DOI: [10.1017/CB09780511804090](https://doi.org/10.1017/CB09780511804090) (cited on pp. 4, 10, 141).
- [5] Andrea ASPERTI, Vincent DANOS, Cosimo LANEVE, and Laurent REGNIER. “Paths in the lambda-calculus”. In: *LICS*. IEEE Computer Society, 1994, pages 426–436. DOI: [10.1109/LICS.1994.316048](https://doi.org/10.1109/LICS.1994.316048) (cited on p. 17).
- [6] Mikhail J. ATALLAH and Marina BLANTON, editors. *Algorithms and Theory of Computation Handbook*. CRC, 1998. DOI: [10.1201/9781584888239](https://doi.org/10.1201/9781584888239) (cited on p. 4).
- [7] Clément AUBERT. “Réseaux de preuves booléens sous-logarithmiques”. Master’s thesis. L.I.P.N.: L.M.F.I., Paris VII, Sept. 2010 (cited on p. 32).
- [8] Clément AUBERT. “Sublogarithmic uniform Boolean proof nets”. In: *DICE*. Edited by Jean-Yves Marion. Volume 75. Electronic Proceedings in Theoretical Computer Science. 2011, pages 15–27. DOI: [10.4204/EPTCS.75.2](https://doi.org/10.4204/EPTCS.75.2) (cited on p. 32).
- [9] Clément AUBERT and Marc BAGNOL. “Unification and Logarithmic Space”. In: *RTA-TLCA*. Edited by Gilles Dowek. Volume 8650. Lecture Notes in Computer Science. Springer, 2014, pages 77–92. DOI: [10.1007/978-3-319-08918-8_6](https://doi.org/10.1007/978-3-319-08918-8_6) (cited on pp. 90, 120).
- [10] Clément AUBERT and Thomas SEILLER. “Characterizing co-NL by a group action”. In: *Arxiv preprint* abs/1209.3422 (2012). arXiv: [1209.3422](https://arxiv.org/abs/1209.3422) [cs.LG] (cited on pp. 58, 69 sq., 82 sq., 94, 106, 110, 112, 141).
- [11] Clément AUBERT and Thomas SEILLER. “Logarithmic Space and Permutations”. In: *Arxiv preprint* abs/1301.3189 (2013). arXiv: [1301.3189](https://arxiv.org/abs/1301.3189) [cs.LG] (cited on pp. 58, 82, 94, 106, 110).
- [12] Sheldon AXLER. *Linear Algebra Done Right*. Undergraduate Texts in Mathematics. Springer, 1997. DOI: [10.1007/b97662](https://doi.org/10.1007/b97662) (cited on p. 125).

- [13] Patrick BAILLOT and Damiano MAZZA. “Linear Logic by Levels and Bounded Time Complexity”. In: *Theoretical Computer Science* 411.2 (2010), pages 470–503. DOI: [10.1016/j.tcs.2009.09.015](https://doi.org/10.1016/j.tcs.2009.09.015) (cited on p. 18).
- [14] Patrick BAILLOT and Marco PEDICINI. “Elementary Complexity and Geometry of Interaction”. In: *Fundamenta Informaticae* 45.1–2 (2001), pages 1–31 (cited on pp. 18, 58).
- [15] David A. BARRINGTON, Neil IMMERMAN, and Howard STRAUBING. “On uniformity within NC¹”. In: *Journal of Computer and System Sciences* 41.3 (1990), pages 274–306. DOI: [10.1016/0022-0000\(90\)90022-D](https://doi.org/10.1016/0022-0000(90)90022-D) (cited on pp. 11, 23).
- [16] Amir M. BEN-AMRAM. “What is a “Pointer Machine”?” In: *Science of Computer Programming*. Volume 26. 2. Association for Computing Machinery, June 1995, pages 88–95. DOI: [10.1145/202840.202846](https://doi.org/10.1145/202840.202846) (cited on p. 94).
- [17] Amir M. BEN-AMRAM and Neil D. JONES. “Computational Complexity via Programming Languages: Constant Factors Do Matter”. In: *Acta Informatica* 37.2 (Oct. 2000), pages 83–120. DOI: [10.1007/s002360000038](https://doi.org/10.1007/s002360000038) (cited on p. 123).
- [18] Guillaume BONFANTE and Virgile MOGBIL. *A circuit uniformity sharper than DLogTime*. Technical report. CARTE - LIPN, Apr. 2012 (cited on p. 23).
- [19] Allan BORODIN. “On Relating Time and Space to Size and Depth”. In: *SIAM Journal on Computing* 6.4 (1977), pages 733–744. DOI: [10.1137/0206054](https://doi.org/10.1137/0206054) (cited on p. 22).
- [20] Allan BORODIN, Stephen Arthur COOK, Patrick W. DYMOND, Walter L. RUZZO, and Martin TOMPA. “Two Applications of Inductive Counting for Complementation Problems”. In: *SIAM Journal on Computing* 18.3 (1989), pages 559–578. DOI: [10.1137/0218038](https://doi.org/10.1137/0218038) (cited on p. 54).
- [21] Egon BÖRGER, editor. *Journal for Universal Computer Science* 3.4 (1997): *Ten years of Gurevich’s Abstract State Machines*. DOI: [10.3217/jucs-003-04](https://doi.org/10.3217/jucs-003-04) (cited on p. 97).
- [22] Jin-Yi CAI, Martin FÜRER, and Neil IMMERMAN. “An optimal lower bound on the number of variables for graph identifications”. In: *Combinatorica* 12.4 (1992), 389–410. DOI: [10.1007/BF01305232](https://doi.org/10.1007/BF01305232) (cited on p. 95).
- [23] Liming CAI and Jianer CHEN. “On input read-modes of alternating Turing machines”. In: *Theoretical Computer Science* 148.1 (1995), pages 33–55. DOI: [10.1016/0304-3975\(94\)00253-F](https://doi.org/10.1016/0304-3975(94)00253-F) (cited on p. 23).
- [24] Ashok K. CHANDRA, Dexter KOZEN, and Larry J. STOCKMEYER. “Alternation”. In: *Journal of the ACM* 28.1 (Jan. 1981), pages 114–133. DOI: [10.1145/322234.322243](https://doi.org/10.1145/322234.322243) (cited on pp. 5 sq., 9).
- [25] Ashok K. CHANDRA, Larry J. STOCKMEYER, and Uzi VISHKIN. “Constant Depth Reducibility”. In: *SIAM Journal on Computing* 13.2 (1984), 423–439. DOI: [10.1137/0213028](https://doi.org/10.1137/0213028) (cited on p. 11).
- [26] John B. CONWAY. *A Course in Functional Analysis*. Volume 96. Graduate Texts in Mathematics. Springer, 1990 (cited on p. 126).
- [27] Stephen Arthur COOK. “A Taxinomy of Problems with Fast Parallel Algorithms”. In: *Information and Control* 64.1–3 (1985), pages 2–22. DOI: [10.1016/S0019-9958\(85\)80041-3](https://doi.org/10.1016/S0019-9958(85)80041-3) (cited on p. 45).

- [28] Stephen Arthur COOK and Pierre MCKENZIE. “Problems Complete for Deterministic Logarithmic Space”. In: *Journal of Algorithms* 8.3 (1987), pages 385–394. DOI: [10.1016/0196-6774\(87\)90018-6](https://doi.org/10.1016/0196-6774(87)90018-6) (cited on p. 12).
- [29] Stephen Arthur COOK and Charles W. RACKOFF. “Space Lower Bounds for Maze Threadability on Restricted Machines”. In: *SIAM Journal on Computing* 9.3 (1980), pages 636–652. DOI: [10.1137/0209048](https://doi.org/10.1137/0209048) (cited on p. 95).
- [30] Pierre-Louis CURIEN. *Notes on operator algebras, part I*. Lecture Notes. Mar. 2011 (cited on p. 125).
- [31] Ugo DAL LAGO. “Context Semantics, Linear Logic, and Computational Complexity”. In: *ACM Transactions on Computational Logic* 10.4 (Aug. 2009), 25:1–25:32. DOI: [10.1145/1555746.1555749](https://doi.org/10.1145/1555746.1555749) (cited on p. 58).
- [32] Ugo DAL LAGO and Claudia FAGGIAN. “On Multiplicative Linear Logic, Modality and Quantum Circuits”. In: *QPL*. Edited by Bart Jacobs, Peter Selinger, and Bas Spitters. Volume 95. Electronic Proceedings in Theoretical Computer Science. 2011, pages 55–66. DOI: [10.4204/EPTCS.95.6](https://doi.org/10.4204/EPTCS.95.6) (cited on p. 59).
- [33] Ugo DAL LAGO and Olivier LAURENT. “Quantitative Game Semantics for Linear Logic”. In: *CSL*. Edited by Michael Kaminski and Simone Martini. Volume 5213. Lecture Notes in Computer Science. Springer, 2008, pages 230–245. DOI: [10.1007/978-3-540-87531-4_18](https://doi.org/10.1007/978-3-540-87531-4_18) (cited on p. 58).
- [34] Ugo DAL LAGO, Simone MARTINI, and Margherita ZORZI. “General Ramified Recurrence is Sound for Polynomial Time”. In: *DICE*. Edited by Patrick Baillot. Volume 23. Electronic Proceedings in Theoretical Computer Science. 2010, pages 47–62. DOI: [10.4204/EPTCS.23.4](https://doi.org/10.4204/EPTCS.23.4) (cited on p. 3).
- [35] Ugo DAL LAGO and Ulrich SCHÖPP. “Type Inference for Sublinear Space Functional Programming”. In: *APLAS*. Edited by Kazunori Ueda. Volume 6461. Lecture Notes in Computer Science. Springer, 2010, pages 376–391. DOI: [10.1007/978-3-642-17164-2_26](https://doi.org/10.1007/978-3-642-17164-2_26) (cited on p. 18).
- [36] Vincent DANOS. “La Logique Linéaire appliquée à l’étude de divers processus de normalisation (principalement du λ -calcul)”. PhD thesis. Université Paris VII, 1990 (cited on pp. 37 sq.).
- [37] Vincent DANOS and Jean-Baptiste JOINET. “Linear Logic & Elementary Time”. In: *Information and Computation* 183.1 (2003), pages 123–137. DOI: [10.1016/S0890-5401\(03\)00010-5](https://doi.org/10.1016/S0890-5401(03)00010-5) (cited on p. 18).
- [38] Vincent DANOS and Laurent REGNIER. “The Structure of Multiplicatives”. In: *Archive for Mathematical logic* 28.3 (1989), pages 181–203. DOI: [10.1007/BF01622878](https://doi.org/10.1007/BF01622878) (cited on pp. 19, 24, 26).
- [39] Daniel DE CARVALHO, Michele PAGANI, and Lorenzo TORTORA DE FALCO. “A Semantic Measure of the Execution Time in Linear Logic”. In: *Theoretical Computer Science* 412.20 (Apr. 2011): *Girard’s Festschrift*. Edited by Thomas Ehrhard, Claudia Faggian, and Olivier Laurent, pages 1884–1902. DOI: [10.1016/j.tcs.2010.12.017](https://doi.org/10.1016/j.tcs.2010.12.017) (cited on p. 18).
- [40] Heinz-Dieter EBBINGHAUS and Jörg FLUM. *Finite model theory*. Perspectives in Mathematical Logic. Springer, 1995, pages I–XV, 1–327 (cited on p. 94).

- [41] Thomas EHRHARD, Claudia FAGGIAN, and Olivier LAURENT, editors. *Theoretical Computer Science* 412.20 (Apr. 2011): *Girard's Festschrift*.
- [42] Thomas EHRHARD, Jean-Yves GIRARD, Paul RUET, and Philip J. SCOTT, editors. *Linear Logic in Computer Science*. London Mathematical Society Lecture Note Series 316. Cambridge University Press, 2004. DOI: [10.1017/CB09780511550850](https://doi.org/10.1017/CB09780511550850).
- [43] Igor FULMAN. *Crossed products of von Neumann algebras by equivalence relations and their subalgebras*. Volume 126. Memoirs of the American Mathematical Society 602. AMS Bookstore, 1997. DOI: [10.1090/memo/0602](https://doi.org/10.1090/memo/0602) (cited on p. 90).
- [44] Merrick L. FURST, James B. SAXE, and Michael SIPSER. “Parity, Circuits, and the Polynomial-Time Hierarchy”. In: *Theory of Computing Systems* 17.1 (1984), pages 13–27. DOI: [10.1007/BF01744431](https://doi.org/10.1007/BF01744431) (cited on p. 11).
- [45] Gerhard GENTZEN. “Untersuchungen über das Logische Schliessen”. In: *Mathematische Zeitschrift* 39 (1935), 176–210 and 405–431 (cited on p. 14).
- [46] Gerhard GENTZEN. *The Collected Papers of Gerhard Gentzen*. Edited by M. E. Szabo. Studies in Logic and the Foundations of Mathematics. North-Holland, 1969 (cited on p. 14).
- [47] Jean-Yves GIRARD. “Linear logic”. In: *Theoretical Computer Science* 50.1 (1987), pages 1–101. DOI: [10.1016/0304-3975\(87\)90045-4](https://doi.org/10.1016/0304-3975(87)90045-4) (cited on pp. 12, 16, 28).
- [48] Jean-Yves GIRARD. “Geometry of interaction I: Interpretation of System F”. In: *Studies in Logic and the Foundations of Mathematics* 127 (1989), pages 221–260. DOI: [10.1016/S0049-237X\(08\)70271-4](https://doi.org/10.1016/S0049-237X(08)70271-4) (cited on pp. 58, 106, 110).
- [49] Jean-Yves GIRARD. “Towards a geometry of interaction”. In: *Proceedings of the AMS-IMS-SIAM Joint Summer Research Conference held June 14-20, 1987*. Edited by John W. Gray and Andre Šcedrov. Volume 92. Categories in Computer Science and Logic. American Mathematical Society, 1989, pages 69–108. DOI: [10.1090/conm/092/1003197](https://doi.org/10.1090/conm/092/1003197) (cited on pp. 16 sq., 19, 54, 63, 106).
- [50] Jean-Yves GIRARD. “Geometry of interaction III: accommodating the additives”. In: *Advances in Linear Logic*. Edited by Jean-Yves Girard, Yves Lafont, and Laurent Regnier. London Mathematical Society Lecture Note Series 222. Cambridge University Press, June 1995, pages 329–389. DOI: [10.1017/CB09780511629150.017](https://doi.org/10.1017/CB09780511629150.017) (cited on pp. 76, 126).
- [51] Jean-Yves GIRARD. “Proof-Nets: The Parallel Syntax for Proof-Theory”. In: *Logic and Algebra*. Lecture Notes in Pure and Applied Mathematics 180 (May 1996), pages 97–124 (cited on p. 14).
- [52] Jean-Yves GIRARD. “Light Linear Logic”. In: *Information and Computation* 143.2 (1998), pages 175–204. DOI: [10.1006/inco.1998.2700](https://doi.org/10.1006/inco.1998.2700) (cited on p. 18).
- [53] Jean-Yves GIRARD. *Introduction aux algèbres d'opérateurs I: Des espaces de Hilbert aux algèbres stellaires*. Lecture Notes. Sept. 2002 (cited on p. 125).
- [54] Jean-Yves GIRARD. “Between Logic and Quantic: a Tract”. In: *Linear Logic in Computer Science*. Edited by Thomas Ehrhard, Jean-Yves Girard, Paul Ruet, and Philip J. Scott. London Mathematical Society Lecture Note Series 316. Cambridge University Press, 2004, pages 346–381. DOI: [10.1017/CB09780511550850.011](https://doi.org/10.1017/CB09780511550850.011) (cited on p. 59).
- [55] Jean-Yves GIRARD. *Le Point Aveugle I: Vers la Perfection*. Visions des sciences. Hermann, July 2006 (cited on p. 12).

- [56] Jean-Yves GIRARD. *Le Point Aveugle II : Vers l'Imperfection*. Visions des sciences. Hermann, Mar. 2007 (cited on pp. 12, 90).
- [57] Jean-Yves GIRARD. "Geometry of Interaction V: logic in the hyperfinite factor". In: *Theoretical Computer Science* 412.20 (Apr. 2011): *Girard's Festschrift*. Edited by Thomas Ehrhard, Claudia Faggian, and Olivier Laurent, pages 1860–1883. DOI: [10.1016/j.tcs.2010.12.016](https://doi.org/10.1016/j.tcs.2010.12.016) (cited on pp. 16, 57, 61).
- [58] Jean-Yves GIRARD. *The Blind Spot: Lectures on Logic*. Edited by European Mathematical Society. Sept. 2011, page 550. DOI: [10.4171/088](https://doi.org/10.4171/088) (cited on pp. 12, 60).
- [59] Jean-Yves GIRARD. "Normativity in Logic". In: *Epistemology versus Ontology*. Edited by Peter Dybjer, Sten Lindström, Erik Palmgren, and Göran Sundholm. Volume 27. Logic, Epistemology, and the Unity of Science. Springer, 2012, pages 243–263. DOI: [10.1007/978-94-007-4435-6_12](https://doi.org/10.1007/978-94-007-4435-6_12) (cited on pp. 58, 61, 67, 80–82, 93, 111).
- [60] Jean-Yves GIRARD. "Three lightings of logic". In: *CSL*. Edited by Simona Ronchi Della Rocca. Volume 23. Leibniz International Proceedings in Informatics. Schloss Dagstuhl - Leibniz-Zentrum für Informatik, 2013, pages 11–23. DOI: [10.4230/LIPIcs.CSL.2013.11](https://doi.org/10.4230/LIPIcs.CSL.2013.11) (cited on p. 90).
- [61] Jean-Yves GIRARD, Yves LAFONT, and Laurent REGNIER, editors. *Advances in Linear Logic*. London Mathematical Society Lecture Note Series 222. Cambridge University Press, June 1995.
- [62] Jean-Yves GIRARD, Yves LAFONT, and Paul TAYLOR. *Proofs and Types*. Volume 7. Cambridge Tracts in Theoretical Computer Science. Cambridge University Press, 1989 (cited on pp. 12, 63, 90).
- [63] Jean-Yves GIRARD, Andre SCEDROV, and Philip J. SCOTT. "Bounded linear logic: a modular approach to polynomial-time computability". In: *Theoretical Computer Science* 97.1 (1992), pages 1–66. DOI: [10.1016/0304-3975\(92\)90386-T](https://doi.org/10.1016/0304-3975(92)90386-T) (cited on p. 18).
- [64] Leslie M. GOLDSCHLAGER. "A Unified Approach to Models of Synchronous Parallel Machines". In: *STOC*. Edited by Richard J. Lipton, Walter A. Burkhard, Walter J. Savitch, Emily P. Friedman, and Alfred V. Aho. Association for Computing Machinery, 1978, pages 89–94. DOI: [10.1145/800133.804336](https://doi.org/10.1145/800133.804336) (cited on p. 5).
- [65] Raymond GREENLAW, H. James HOOVER, and Walter L. RUZZO. *Limits to Parallel Computation : P-Completeness theory*. Oxford UP, 1995 (cited on pp. 4, 22).
- [66] Erich GRÄDEL and Gregory L. MCCOLM. "On the Power of Deterministic Transitive Closures". In: *Information and Computation* 119.1 (1995), 129–135. DOI: [10.1006/inco.1995.1081](https://doi.org/10.1006/inco.1995.1081) (cited on p. 95).
- [67] Stefano GUERRINI and Andrea MASINI. "Parsing MELL proof nets". In: *Theoretical Computer Science* 254.1–2 (2001), pages 317–335. DOI: [10.1016/S0304-3975\(99\)00299-6](https://doi.org/10.1016/S0304-3975(99)00299-6) (cited on p. 37).
- [68] Uffe HAAGERUP. "The Standard Form of von Neumann Algebras". In: *Mathematica Scandinavica* 37.271–283 (1975) (cited on p. 137).
- [69] Juris HARTMANIS. "On Non-Determinacy in Simple Computing Devices". In: *Acta Informatica* 1.4 (1972), pages 336–344. DOI: [10.1007/BF00289513](https://doi.org/10.1007/BF00289513) (cited on pp. 98 sq., 104).

- [70] William HESSE, Eric W. ALLENDER, and David A. BARRINGTON. “Uniform constant-depth threshold circuits for division and iterated multiplication”. In: *Journal of Computer and System Sciences* 65.4 (2002), pages 695–716. DOI: [10.1016/S0022-0000\(02\)00025-9](https://doi.org/10.1016/S0022-0000(02)00025-9) (cited on p. 23).
- [71] Martin HOFMANN, Ramyaa RAMYAA, and Ulrich SCHÖPP. “Pure Pointer Programs and Tree Isomorphism”. In: *FoSSaCS*. Edited by Frank Pfenning. Volume 7794. Lecture Notes in Computer Science. Springer, 2013, pages 321–336. DOI: [10.1007/978-3-642-37075-5_21](https://doi.org/10.1007/978-3-642-37075-5_21) (cited on p. 95).
- [72] Martin HOFMANN and Ulrich SCHÖPP. “Pointer Programs and Undirected Reachability”. In: *LICS*. IEEE Computer Society, 2009, pages 133–142. DOI: [10.1109/LICS.2009.41](https://doi.org/10.1109/LICS.2009.41) (cited on p. 95).
- [73] Martin HOFMANN and Ulrich SCHÖPP. “Pure Pointer Programs with Iteration”. In: *ACM Transactions on Computational Logic* 11.4 (2010). DOI: [10.1145/1805950.1805956](https://doi.org/10.1145/1805950.1805956) (cited on p. 95).
- [74] Markus HOLZER and Martin KUTRIB. “Descriptive and computational complexity of finite automata - A survey”. In: *Information and Computation* 209.3 (2011), pages 456–470. DOI: [10.1016/j.ic.2010.11.013](https://doi.org/10.1016/j.ic.2010.11.013) (cited on p. 106).
- [75] Markus HOLZER, Martin KUTRIB, and Andreas MALCHER. “Multi-Head Finite Automata: Characterizations, Concepts and Open Problems”. In: *CSP*. Edited by Turlough Neary, Damien Woods, Anthony Karel Seda, and Niall Murphy. Volume 1. Electronic Proceedings in Theoretical Computer Science. 2008, pages 93–107. DOI: [10.4204/EPTCS.1.9](https://doi.org/10.4204/EPTCS.1.9) (cited on pp. 98, 104).
- [76] John E. HOPCROFT, Rajeev MOTWANI, and Jeffrey D. ULLMAN. *Introduction to automata theory, languages, and computation*. international edition (2. ed). Addison-Wesley, 2003, pages I–XIV, 1–521 (cited on pp. 98, 107).
- [77] Neil IMMERMAN. “Languages that Capture Complexity Classes”. In: *SIAM Journal on Computing* 16.4 (1987), pages 760–778. DOI: [10.1137/0216051](https://doi.org/10.1137/0216051) (cited on p. 3).
- [78] Neil IMMERMAN. “Nondeterministic space is closed under complementation”. In: *CoCo*. IEEE Computer Society, 1988, pages 112–115. DOI: [10.1109/SCT.1988.5270](https://doi.org/10.1109/SCT.1988.5270) (cited on p. 11).
- [79] Neil IMMERMAN. *Descriptive complexity*. Graduate Texts in Computer Science. Springer, 1999, pages I–XVI, 1–268. DOI: [10.1007/978-1-4612-0539-5](https://doi.org/10.1007/978-1-4612-0539-5) (cited on p. 94).
- [80] Paulin JACOBÉ DE NAUROIS and Virgile MOGBIL. “Correctness of linear logic proof structures is NL-complete”. In: *Theoretical Computer Science* 412.20 (Apr. 2011): *Girard’s Festschrift*. Edited by Thomas Ehrhard, Claudia Faggian, and Olivier Laurent, pages 1941–1957. DOI: [10.1016/j.tcs.2010.12.020](https://doi.org/10.1016/j.tcs.2010.12.020) (cited on p. 17).
- [81] Birgit JENNER, Klaus-Jörn LANGE, and Pierre MCKENZIE. *Tree Isomorphism and Some Other Complete Problems for Deterministic Logspace*. Technical report. Université de Montréal - DIRO, 1997 (cited on p. 11).
- [82] Tao JIANG, Ming LI, and Bala RAVIKUMAR. “Basic notions in computational complexity”. In: *Algorithms and Theory of Computation Handbook*. Edited by Mikhail J. Atallah and Marina Blanton. CRC, 1998. DOI: [10.1201/9781584888239](https://doi.org/10.1201/9781584888239) (cited on p. 4).

- [83] Jean-Baptiste JOINET. “Logique et interaction”. Habilitation à diriger des recherches. Université Paris Diderot (Paris 7), Dec. 2007 (cited on pp. 12, 15).
- [84] Neil D. JONES. “Constant Time Factors *Do* Matter”. In: *STOC*. Edited by S. Rao Kosaraju, David S. Johnson, and Alok Aggarwal. San Diego, California, USA: Association for Computing Machinery, 1993, pages 602–611. DOI: [10.1145/167088.167244](https://doi.org/10.1145/167088.167244) (cited on p. 123).
- [85] Neil D. JONES, Y. Edmund LIEN, and William T. LAASER. “New Problems Complete for Nondeterministic Log Space”. In: *Mathematical Systems Theory* 10.1 (1976), pages 1–17. DOI: [10.1007/BF01683259](https://doi.org/10.1007/BF01683259) (cited on p. 11).
- [86] Richard V. KADISON and John R. RINGROSE. *Fundamentals of the theory of operator algebras. Vol. I*. Volume 15. Graduate Studies in Mathematics. 1997 (cited on p. 125).
- [87] Max I. KANOVICH. “Horn Programming in Linear Logic Is NP-Complete”. In: *LICS*. IEEE Computer Society, 1992, 200–210. DOI: [10.1109/LICS.1992.185533](https://doi.org/10.1109/LICS.1992.185533) (cited on p. 17).
- [88] K. N. KING. “Alternating Multihead Finite Automata”. In: *Theoretical Computer Science* 61.2–3 (1988), pages 149–174. DOI: [10.1016/0304-3975\(88\)90122-3](https://doi.org/10.1016/0304-3975(88)90122-3) (cited on p. 106).
- [89] Andrey Nikolaevich KOLMOGOROV and Vladimir Andreyevich USPENSKIĬ. “On the definition of an algorithm”. In: *Uspekhi Matematicheskikh Nauk* 13.4(82) (1958), 3–28 (cited on p. 96).
- [90] Andrey Nikolaevich KOLMOGOROV and Vladimir Andreyevich USPENSKIĬ. “On the definition of an algorithm”. In: *American Mathematical Society Translations, Series 2* 29 (1963), 217–245 (cited on p. 96).
- [91] Dexter KOZEN. “On Two Letters versus Three”. In: *FICS*. Edited by Zoltán Ésik and Anna Ingólfssdóttir. Volume NS-02-2. BRICS Notes Series. University of Aarhus, July 2002, pages 44–50 (cited on p. 100).
- [92] Jean-Louis KRIVINE. *Lambda-calculus, types and models*. Ellis Horwood series in computers and their applications. Masson, 1993, pages I–VIII, 1–180 (cited on p. 14).
- [93] Yves LAFONT. “From Proof-Nets to Interaction Nets”. In: *Advances in Linear Logic*. Edited by Jean-Yves Girard, Yves Lafont, and Laurent Regnier. London Mathematical Society Lecture Note Series 222. Cambridge University Press, June 1995, pages 225–247. DOI: [10.1017/CB09780511629150.012](https://doi.org/10.1017/CB09780511629150.012) (cited on p. 37).
- [94] Joachim LAMBEK. “The Mathematics of Sentence Structure”. In: *The American Mathematical Monthly* 65.3 (1958), 154–170. DOI: [10.2307/2310058](https://doi.org/10.2307/2310058) (cited on p. 12).
- [95] Olivier LAURENT. *Théorie de la démonstration*. LMFI’s master *Polycopié* (cited on p. 13).
- [96] Olivier LAURENT. “A Token Machine for Full Geometry of Interaction (Extended Abstract)”. In: *Typed Lambda Calculi and Applications*. Edited by Samson Abramsky. Volume 2044. Lecture Notes in Computer Science. Springer, May 2001, pages 283–297. DOI: [10.1007/3-540-45413-6_23](https://doi.org/10.1007/3-540-45413-6_23) (cited on p. 17).
- [97] Daniel LEIVANT and Jean-Yves MARION. “Evolving Graph-Structures and Their Implicit Computational Complexity”. In: *ICALP(2)*. Edited by Fedor V. Fomin, Rusins Martins Freivalds, Marta Z. Kwiatkowska, and David Peleg. Volume 7966. Lecture Notes in Computer Science. Springer, July 2013, 349–360. DOI: [10.1007/978-3-642-39212-2_32](https://doi.org/10.1007/978-3-642-39212-2_32) (cited on p. 97).

- [98] Patrick D. LINCOLN. “Deciding Provability of Linear Logic Formulas”. In: *London Mathematical Society Lecture Note Series* (1995), 109–122. DOI: [10.1017/CB09780511629150.006](https://doi.org/10.1017/CB09780511629150.006) (cited on p. 17).
- [99] Philip J. MAHER. “Some operator inequalities concerning generalized inverses”. In: *Illinois Journal of Mathematics* 34.3 (1990), 503–514 (cited on p. 118).
- [100] Harry MAIRSON and Kazushige TERUI. “On the Computational Complexity of Cut-Elimination in Linear Logic”. In: *Theoretical Computer Science* (2003), pages 23–36. DOI: [10.1007/978-3-540-45208-9_4](https://doi.org/10.1007/978-3-540-45208-9_4) (cited on pp. 17, 50, 59).
- [101] François MAUREL. “Nondeterministic Light Logics and NP-Time”. In: *Typed Lambda Calculi and Applications*. Edited by Martin Hofmann. Volume 2701. Lecture Notes in Computer Science. Springer, 2003, pages 241–255. DOI: [10.1007/3-540-44904-3_17](https://doi.org/10.1007/3-540-44904-3_17) (cited on p. 54).
- [102] Damiano MAZZA. “Non-uniform Polytime Computation in the Infinitary Affine Lambda-Calculus”. In: *ICALP (2)*. Edited by Javier Esparza, Pierre Fraigniaud, Thore Husfeldt, and Elias Koutsoupias. Volume 8573. Lecture Notes in Computer Science. Springer, 2014, pages 305–317. DOI: [0.1007/978-3-662-43951-7_26](https://doi.org/0.1007/978-3-662-43951-7_26) (cited on p. 55).
- [103] Paul-André MELLIÈS. “A Topological Correctness Criterion for Multiplicative Non-Commutative Logic”. In: *Linear Logic in Computer Science*. Edited by Thomas Ehrhard, Jean-Yves Girard, Paul Ruet, and Philip J. Scott. London Mathematical Society Lecture Note Series 316. Cambridge University Press, 2004, pages 283–322. DOI: [10.1017/CB09780511550850.009](https://doi.org/10.1017/CB09780511550850.009) (cited on p. 54).
- [104] Paul-André MELLIÈS and Nicolas TABAREAU. “Resource modalities in tensor logic”. In: *Annals of Pure and Applied Logic* 161.5 (2010), pages 632–653. DOI: [10.1016/j.apal.2009.07.018](https://doi.org/10.1016/j.apal.2009.07.018) (cited on p. 123).
- [105] Virgile MOGBIL. “Non-deterministic Boolean Proof Nets”. In: *FOPARA*. Edited by Marko C. J. D. van Eekelen and Olha Shkaravska. Volume 6324. Lecture Notes in Computer Science. Springer, 2009, pages 131–145. DOI: [10.1007/978-3-642-15331-0_9](https://doi.org/10.1007/978-3-642-15331-0_9) (cited on pp. 22, 32, 54).
- [106] Virgile MOGBIL. “Complexité en logique linéaire, et logique linéaire en complexité implicite”. Habilitation à diriger des recherches en Informatique. Université Paris XIII, Nov. 2012 (cited on pp. 13–15).
- [107] Virgile MOGBIL and Vincent RAHLI. “Uniform Circuits, & Boolean Proof Nets”. In: *LFCs*. Edited by Sergei N. Artëmov and Anil Nerode. Volume 4514. Lecture Notes in Computer Science. Springer, 2007, pages 401–421. DOI: [10.1007/978-3-540-72734-7_28](https://doi.org/10.1007/978-3-540-72734-7_28) (cited on pp. 22, 32, 40, 43 sq.).
- [108] Burkhard MONIEN. “Two-Way Multihead Automata Over a One-Letter Alphabet”. In: *RAIRO - Theoretical Informatics and Applications - Informatique Théorique et Applications* 14.1 (1980), pages 67–82 (cited on p. 104).
- [109] Gerard J. MURPHY. *C*-algebras and operator theory*. Academic Press, 1990 (cited on pp. 125 sq., 132).
- [110] Michael Aaron NIELSEN and Isaac L. CHUANG. *Quantum computation and quantum information*. Cambridge University Press, 2010. DOI: [10.1017/CB09780511976667](https://doi.org/10.1017/CB09780511976667) (cited on p. 59).

- [III] Christos H. PAPADIMITRIOU. *Computational Complexity*. Mathematics / a second level course v. 3-4. Addison-Wesley, 1995 (cited on pp. 4 sq.).
- [II2] Bruno POIZAT. *Les petits cailloux*. Aléas, 1995 (cited on p. 22).
- [II3] Kenneth W. REGAN and Heribert VOLLMER. “Gap-Languages and Log-Time Complexity Classes”. In: *Theoretical Computer Science* 188 (1997), pages 101–116. DOI: [10.1016/S0304-3975\(96\)00288-5](https://doi.org/10.1016/S0304-3975(96)00288-5) (cited on p. 23).
- [II4] Omer REINGOLD. “Undirected connectivity in log-space”. In: *Journal of the ACM* 55.4 (Sept. 2008), 17:1–17:24. DOI: [10.1145/1391289.1391291](https://doi.org/10.1145/1391289.1391291) (cited on p. 12).
- [II5] Walter L. RUZZO. “On Uniform Circuit Complexity”. In: *Journal of Computer and System Sciences* 22.3 (June 1981), pages 365–383. DOI: [10.1016/0022-0000\(81\)90038-6](https://doi.org/10.1016/0022-0000(81)90038-6) (cited on pp. 11, 23).
- [II6] Arnold SCHÖNHAGE. “Storage Modification Machines”. In: *SIAM Journal on Computing* 9.3 (1980), pages 490–508. DOI: [10.1137/0209036](https://doi.org/10.1137/0209036) (cited on pp. 96 sq.).
- [II7] Ulrich SCHÖPP. “Stratified Bounded Affine Logic for Logarithmic Space”. In: *LICS*. IEEE Computer Society, 2007, pages 411–420. DOI: [10.1109/LICS.2007.45](https://doi.org/10.1109/LICS.2007.45) (cited on p. 18).
- [II8] Dana S. SCOTT. “Domains for Denotational Semantics”. In: *ICALP*. Edited by Mogens Nielsen and Erik Meineche Schmidt. Volume 140. Lecture Notes in Computer Science. Springer, 1982, 577–613. DOI: [10.1007/BFb0012801](https://doi.org/10.1007/BFb0012801) (cited on p. 16).
- [II9] Thomas SEILLER. “Interaction Graphs: Additives”. In: *Arxiv preprint abs/1205.6557* (2012). arXiv: [1205.6557](https://arxiv.org/abs/1205.6557) [cs.LG] (cited on p. 69).
- [I20] Thomas SEILLER. “Interaction Graphs: Multiplicatives”. In: *Annals of Pure and Applied Logic* 163 (Dec. 2012), pages 1808–1837. DOI: [10.1016/j.apal.2012.04.005](https://doi.org/10.1016/j.apal.2012.04.005) (cited on pp. 69, 81, 126).
- [I21] Thomas SEILLER. “Logique dans le facteur hyperfini : géometrie de l’interaction et complexité”. PhD thesis. Université de la Méditerranée, 2012 (cited on pp. 64 sq., 69, 81–83, 123, 126).
- [I22] Peter SELINGER. “Towards a quantum programming language”. In: *Mathematical Structures in Computer Science* 14.4 (2004), pages 527–586. DOI: [10.1017/S0960129504004256](https://doi.org/10.1017/S0960129504004256) (cited on p. 59).
- [I23] Michael SIPSER. “Borel Sets and Circuit Complexity”. In: *STOC*. Edited by David S. Johnson, Ronald Fagin, Michael L. Fredman, David Harel, Richard M. Karp, Nancy A. Lynch, Christos H. Papadimitriou, Ronald L. Rivest, Walter L. Ruzzo, and Joel I. Seiferas. Association for Computing Machinery, 1983, 61–69. DOI: [10.1145/800061.808733](https://doi.org/10.1145/800061.808733) (cited on p. 11).
- [I24] Michael SIPSER. *Introduction to the Theory of Computation*. 3rd edition. Cengage Learning, 2012, pages I–XXII, 1–458 (cited on p. 5).
- [I25] Ivan Hal SUDBOROUGH. “On Tape-Bounded Complexity Classes and Multi-Head Finite Automata”. In: *SWAT(FOCS)*. IEEE Computer Society, 1973, 138–144. DOI: [10.1109/SWAT.1973.20](https://doi.org/10.1109/SWAT.1973.20) (cited on p. 104).
- [I26] Masamichi TAKESAKI. *Theory of Operator Algebras 1*. Volume 124. Encyclopedia of Mathematical Sciences. Springer, 2001 (cited on pp. 126, 132).
- [I27] Masamichi TAKESAKI. *Theory of Operator Algebras 2*. Volume 125. Encyclopedia of Mathematical Sciences. Springer, 2003 (cited on pp. 126, 137).

- [128] Masamichi TAKESAKI. *Theory of Operator Algebras 3*. Volume 127. Encyclopedia of Mathematical Sciences. Springer, 2003 (cited on p. 126).
- [129] Robert Endre TARJAN. “Reference Machines Require Non-linear Time to Maintain Disjoint Sets”. In: *STOC*. Edited by John E. Hopcroft, Emily P. Friedman, and Michael A. Harrison. Association for Computing Machinery, 1977, pages 18–29. DOI: [10.1145/800105.803392](https://doi.org/10.1145/800105.803392) (cited on p. 94).
- [130] Kazushige TERUI. “Proof Nets and Boolean Circuits”. In: *LICS*. IEEE Computer Society, 2004, pages 182–191. DOI: [10.1109/LICS.2004.1319612](https://doi.org/10.1109/LICS.2004.1319612) (cited on pp. 21, 28, 43).
- [131] Peter VAN EMDE BOAS. “Space Measures for Storage Modification Machines”. In: *Information Processing Letters* 30.2 (1989), 103–110. DOI: [10.1016/0020-0190\(89\)90117-8](https://doi.org/10.1016/0020-0190(89)90117-8) (cited on p. 97).
- [132] Peter VAN EMDE BOAS. “Machine Models and Simulation”. In: *Handbook of Theoretical Computer Science. volume A: Algorithms and Complexity*. Edited by Jan Van Leeuwen. Elsevier, 1990, pages 1–66 (cited on p. 4).
- [133] André VAN TONDERVAN. “A Lambda Calculus for Quantum Computation”. In: *SIAM Journal on Computing* 33.5 (2004), 1109–1135. DOI: [10.1137/S0097539703432165](https://doi.org/10.1137/S0097539703432165) (cited on p. 59).
- [134] Hari VENKATESWARAN. “Properties that Characterize LOGCFL”. In: *Journal of Computer and System Sciences* 43.2 (1991), pages 380–404. DOI: [10.1016/0022-0000\(91\)90020-6](https://doi.org/10.1016/0022-0000(91)90020-6) (cited on p. 55).
- [135] Heribert VOLLMER. *Introduction to Circuit Complexity: A Uniform Approach*. Texts in Theoretical Computer Science An EATCS Series. Springer, 1999. DOI: [10.1007/978-3-662-03927-4](https://doi.org/10.1007/978-3-662-03927-4) (cited on pp. iii, 6, 8, 11, 22, 46 sq.).
- [136] Avi WIGDERSON. “The Complexity of Graph Connectivity”. In: *MFCS*. Edited by Ivan M. Havel and Václav Koubek. Volume 629. Lecture Notes in Computer Science. Springer, 1992, 112–132. DOI: [10.1007/3-540-55808-X_10](https://doi.org/10.1007/3-540-55808-X_10) (cited on pp. 12, 44).
- [137] Andrew Chi-Chih YAO and Ronald L. RIVEST. “ $k + 1$ Heads Are Better Than k ”. In: *Journal of the ACM* 25.2 (1978), 337–340. DOI: [10.1145/322063.322076](https://doi.org/10.1145/322063.322076) (cited on p. 104).

Linear Logic and Sub-polynomial Classes of Complexity

Abstract This research in Theoretical Computer Science extends the gateways between Linear Logic and Complexity Theory by introducing two innovative models of computation. It focuses on sub-polynomial classes of complexity: AC and NC —the classes of efficiently parallelizable problems— and L and NL —the deterministic and non-deterministic classes of problems efficiently solvable with low resources on space. Linear Logic is used through its Proof Net presentation to mimic with efficiency the parallel computation of Boolean Circuits, including but not restricted to their constant-depth variants. In a second moment, we investigate how operators on a von Neumann algebra can be used to model computation, thanks to the method provided by the Geometry of Interaction, a subtle reconstruction of Linear Logic. This characterization of computation in logarithmic space with matrices can naturally be understood as a wander on simple structures using pointers, parsing them without modifying them. We make this intuition formal by introducing Non Deterministic Pointer Machines and relating them to other well-known pointer-like-machines. We obtain by doing so new implicit characterizations of sub-polynomial classes of complexity.

Keywords Linear Logic, Complexity, Geometry of Interaction, Operator Algebra, Alternating Turing Machines, Boolean Circuits, Uniformity, Proof Nets.

Discipline Computer science

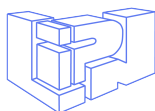


Logique linéaire et classes de complexité sous-polynomiales

Résumé Cette recherche en informatique théorique construit de nouveaux ponts entre logique linéaire et théorie de la complexité. Elle propose deux modèles de machines abstraites qui permettent de capturer de nouvelles classes de complexité avec la logique linéaire, les classes des problèmes efficacement parallélisables (NC et AC) et celle des problèmes solutionnables avec peu d'espace, dans ses versions déterministes et non-déterministes (L et NL). La représentation des preuves de la logique linéaire comme réseaux de preuves est employée pour représenter efficacement le calcul parallèle des circuits booléens, y compris à profondeur constante. La seconde étude s'inspire de la géométrie de l'interaction, une délicate reconstruction de la logique linéaire à l'aide d'opérateurs d'une algèbre de von Neumann. Nous détaillons comment l'interaction d'opérateurs représentant des entiers et d'opérateurs représentant des programmes peut être reconnue nilpotente en espace logarithmique. Nous montrons ensuite comment leur itération représente un calcul effectué par des machines à pointeurs que nous définissons et que nous rattachons à d'autres modèles plus classiques. Ces deux études permettent de capturer de façon implicite de nouvelles classes de complexité, en dessous du temps polynomial.

Mot-clés Logique Linéaire, Complexité, Géométrie de l'interaction, Algèbre d'opérateurs, Machines de Turing Alternantes, Circuits Booléens, Uniformité, Réseaux de Preuves.

Discipline Informatique



Laboratoire d'Informatique de Paris Nord
Institut Galilée - Université Paris-Nord
99, avenue Jean-Baptiste Clément
93430 Villetaneuse, France