



Transistor Level Automatic Generation of Radiation-Hardened Circuits

C. Lazzari

► To cite this version:

C. Lazzari. Transistor Level Automatic Generation of Radiation-Hardened Circuits. Micro and nanotechnologies/Microelectronics. Institut National Polytechnique de Grenoble - INPG, 2007. English. NNT: . tel-00198470

HAL Id: tel-00198470

<https://theses.hal.science/tel-00198470>

Submitted on 17 Dec 2007

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.

INSTITUT NATIONAL POLYTECHNIQUE DE GRENOBLE

N° attribué par la bibliothèque

[illegible]

THESE EN COTUTELLE INTERNATIONALE

pour obtenir le grade de

**DOCTEUR DE L'INP Grenoble
et
de l'Université Fédérale du Rio Grande do Sul**

Spécialité : « Micro et Nano Electronique »

préparée au laboratoire TIMA

dans le cadre de l'Ecole Doctorale « *Electronique, Electrotechnique, Automatique, Télécommunications, Signal* »

présentée et soutenue publiquement

par

Cristiano LAZZARI

le 7 decembre 2007

Transistor Level Automatic Generation of Radiation-Hardened Circuits

Lorena ANGHEL
Ricardo REIS

JURY

M. Raoul Velazco,
M. Matteo SONZA REORDA,
M. Pascal FOUILLAT,
Mme. Lorena ANGHEL,
M. Ricardo REIS,
M. Marcelo LUBASZEWSKI,

Président
Rapporteur
Rapporteur
Directeur de thèse
Co-encadrant
Examineur

CONTENTS

LIST OF ABBREVIATIONS AND ACRONYMS	5
LIST OF FIGURES	7
LIST OF TABLES	11
LIST OF ALGORITHMS	13
ABSTRACT	15
1 INTRODUCTION	17
1.1 Organization of this Thesis	21
2 A TIMING CLOSURE DESIGN FLOW USING A TRANSISTOR LEVEL AUTOMATIC LAYOUT GENERATOR	23
2.1 Introduction	23
2.2 The <i>Super Library</i> Generation	24
2.2.1 The Development of a <i>superlib</i>	26
2.3 The Transistor Level Design Flow	29
2.3.1 The Transistor Level Optimization	31
2.3.2 Transistor Level Optimization for Leakage Reduction	35
2.3.3 The Transistor Level Layout Generation	38
2.4 Transistor-Level vs Traditional Design Flow Summary	40
2.4.1 A Comparison with the Traditional Standard Cell Design Flow	40
2.5 Conclusion	42
3 ALGORITHMS FOR TRANSISTOR LEVEL AUTOMATIC LAYOUT GENERATION	45
3.1 Introduction	45
3.2 An Overview of the Algorithm Classes	45
3.2.1 Deterministic Algorithms	46
3.2.2 Stochastic Algorithms	48
3.3 Goals of Placement Algorithms	49
3.4 Goals of Routing Algorithms	50
3.5 State-of-the-art Algorithms for transistor placement and routing	50
3.5.1 Transistor Placement Using an O-tree Algorithm	50

3.5.2	Transistor Placement with Symmetry Constraints	53
3.5.3	Non-complementary Transistor Placement	54
3.5.4	Transistor Placement & Routing By Using Linear Integer Programming	55
3.5.5	A Maze Routing Steiner Tree with Effective Critical Sink Optimization	56
3.5.6	Routing with a Negotiation-based Algorithm	57
3.6	Transistor Placement Technique Using Genetic Algorithm And Analytical Programming	58
3.6.1	Initial Population Generation	59
3.6.2	The Placement Parameters	59
3.6.3	The Mathematical Modeling	61
3.6.4	Obtained Results	64
3.7	Compaction and Layout Optimization using Linear Programming	65
3.8	Overview of the Algorithms Developed in the Punch++	68
3.9	Conclusion	69
4	A RADIATION-INDUCED EFFECTS OVERVIEW	71
4.1	Introduction	71
4.2	Single Event Effects (SEEs)	72
4.2.1	Soft SEE	73
4.2.2	Single Hard Errors (SHE)	77
4.3	Other Radiation-Induced Effects	78
4.4	Conclusion	78
5	STATE-OF-THE-ART TECHNIQUES FOR SOFT ERROR PROTECTION	79
5.1	Introduction	79
5.2	The Classic Techniques	80
5.3	Gate Duplication Methods	81
5.4	Gate Sizing Techniques	82
5.5	Protecting Sequential Elements Through Feedback Control	82
5.6	Using Time Redundancy To Protect Sequential Elements	83
5.7	Conclusion	86
6	AN EFFICIENT TRANSISTOR SIZING METHODOLOGY FOR SOFT ERROR PROTECTION IN COMBINATIONAL LOGIC CIRCUITS	89
6.1	Introduction	89
6.2	Combinational Circuits Sensitivity	89
6.2.1	The Logical Masking	90
6.2.2	The Electrical Masking	91
6.3	An Analytical Model for Single Event Transients	91
6.3.1	Modeling Resistances and Capacitances	92
6.3.2	The Single Event Transients Model	93
6.3.3	Single Event Transient Propagation	94
6.4	The Transistor Sizing Strategy	95
6.5	The Transistor Sizing Model	96
6.6	Results	98
6.7	Conclusions	100

7 CONCLUSION	103
REFERENCES	105

LIST OF ABBREVIATIONS AND ACRONYMS

ASIC	Application Specific Integrated Circuit
CAD	Computer Aided Design
CTS	Clock Tree Synthesis
CMOS	Complementary Metal-Oxide-Semiconductor
CLIP	Cell Layout via Integer Programming
CVSL	Cascode Voltage Switch Logic
CWSP	Code Word State Preserving
DD	Displacement Damage
DICE	Dual Interlocked storage Cell
DSM	Deep Submicron
EDA	Electronic Design Automation
EDD	Energy Times Delay Squared
GDS	Graphic Data System format
GDSII	Graphic Data System format, second version
GND	Ground
ILP	Integer Linear Programming
IP	Integer Programming
LET	Linear Energy Transfer
LP	Linear Programming
PI	Primary Inputs
PO	Primary Outputs
PTL	Pass Transistor Logic
QP	Quadratic programming
RTL	Register Transfer Language

SAA	South Atlantic Anomaly
SCCG	Static CMOS Complex Gates
SEB	Single Event Burnout
SEE	Single Event Effects
SEGR	Single Event Gate Rupture
SEL	Single Event Latchup
SET	Single Event Transient
SEU	Single Event Upset
SHE	Single Hard Errors
SRAM	Static Random Access Memory
TID	Total Ionizing Dose
TMR	Triple Modular Redundancy
V_{DD}	Supply voltage
VLSI	Very Large Scale Integration

LIST OF FIGURES

Figure 1.1:	Gate and interconnect delay versus technology generation [SIA97]. Delays for feature size below $100nm$ is estimated.	17
Figure 1.2:	Active and static power in microprocessors [Moo03].	18
Figure 1.3:	Clock frequency of high performance ASIC and custom processors in different process technologies [CK02].	19
Figure 2.1:	The proposed design flow overview.	24
Figure 2.2:	Runtime and occupied memory as a function of the number of gates in a library [CAD07a].	26
Figure 2.3:	Number of different complex logic functions in circuits in a commercial $0.35\mu m$ technology. Results were obtained with Cadence RTL Compiler [CAD07b].	27
Figure 2.4:	Generation of the liberty file. This file is used later in the logic synthesis.	30
Figure 2.5:	The proposed transistor level design flow. A) Logic synthesis with the <i>superlib</i> . B) Transistor level optimization tool. C) Layout generation tool, D) Developed tools to cope with the proposed transistor level design flow.	31
Figure 2.6:	Delay and power consumption results for some circuits from the IS-CAS'85 benchmarks.	33
Figure 2.7:	Delay paths of the circuit c1355.	34
Figure 2.8:	An NMOS two-stack inverter reduces the subthreshold leakage when input stays at logic "0".	35
Figure 2.9:	An example of fine-grained sleep transistors.	36
Figure 2.10:	Leakage current in deep submicron technology process.	37
Figure 2.11:	Delay paths distribution before and after transistors length biasing to the benchmarks ISCAS'85.	38
Figure 2.12:	Layout style of two transistor level automatic layout generators. The <i>Punch</i> and <i>Punch++</i> tools.	40
Figure 2.13:	Some layouts of circuits generated by the proposed transistor level design flow.	43
Figure 3.1:	A force-directed example.	46
Figure 3.2:	An Euler path example.	47
Figure 3.3:	Example of a 8-node tree.	51
Figure 3.4:	The O-tree representation placement.	51

Figure 3.5:	Admissible o-tree.	52
Figure 3.6:	Possible insertion position of a external node.	52
Figure 3.7:	Datapath tile placement.	53
Figure 3.8:	Placement with symmetry group and different encodings.	54
Figure 3.9:	Stages of the layout synthesis proposed in [RS03].	55
Figure 3.10:	The CLIP layout style.	55
Figure 3.11:	Comparison of the best trees generated by (a) AMAZE [HNJR07], (b) AHHK [AHH ⁺ 95] and (c) P-Trees [LCLH96] algorithms.	56
Figure 3.12:	Data structure for the negotiation-based algorithm.	57
Figure 3.13:	Euler path example.	59
Figure 3.14:	Transistor orientation constraints.	60
Figure 3.15:	Transistor behavior constraints.	60
Figure 3.16:	Width and height transistors parameters.	61
Figure 3.17:	A row-based boundary representation.	62
Figure 3.18:	Horizontal neighborhood.	63
Figure 3.19:	The Δ_{n_i} representation.	64
Figure 3.20:	Preliminary placement examples.	65
Figure 3.21:	Example of layout.	66
Figure 3.22:	An example of layout compaction. The first figure is a layout com- pacted without cost assignment. Second layout is compacted with costs.	69
Figure 4.1:	Location where Single Event Upsets (SEU) occurred in a spacecraft into a polar orbit of altitude 700km [HSDU ⁺ 90].	71
Figure 4.2:	Classification of radiation-induced effects [Bas06].	72
Figure 4.3:	Heavy ions and protons striking the silicon device. (a) heavy ion increasing the depletion region (b) <i>Spallation</i> caused by a proton or neutron.	73
Figure 4.4:	The current curve as result of a α -particle striking a device according [Mes82].	74
Figure 4.5:	The propagation of a transient fault as results of a particle striking a node of a combinational block.	76
Figure 4.6:	Example of a bit flip in a classic latch.	77
Figure 5.1:	The classic TMR.	80
Figure 5.2:	Two TMR versions with delayed clock and delayed inputs to avoid transient faults coming from the combinational blocks.	80
Figure 5.3:	Latch using the DICE technique as proposed in [CNV96].	82
Figure 5.4:	INV, NOR2 and NAND2 gates using the CWSP technique proposed in [Ang00].	83
Figure 5.5:	Perturbation tolerant circuit based on time redundancy.	84
Figure 5.6:	Using CWSP logic inside the latch as proposed in [LAR05a].	86
Figure 6.1:	The logical masking.	91
Figure 6.2:	The electrical masking.	91

Figure 6.3:	Equivalent circuit for calculating circuit response to an energetic particle hit.	92
Figure 6.4:	A transistor modeled as a resistance.	93
Figure 6.5:	A transient pulse propagation example.	97
Figure 6.6:	Timing penalty versus area overhead for TMR, CWSP and the proposed sizing technique.	100

LIST OF TABLES

Table 2.1:	Maximum number of logic functions obtained by stacked transistors [DRSVW87].	25
Table 2.2:	Synthesis runtime of the ISCAS85 benchmarks with a commercial $0.35\mu m$ standard cell library and the <i>superlib</i> with Cadence RTL Compiler [CAD07b].	27
Table 2.3:	The gate length biasing technique for some ISCAS'85 benchmarks with a $65nm$ technology process.	37
Table 2.4:	Summary of differences of between a traditional standard cell and the proposed transistor level design flow.	41
Table 2.5:	Comparison between layouts generated by the standard cell approach and the proposed transistor level design flow.	42
Table 3.1:	Parameters to the transistors relationship.	61
Table 3.2:	Horizontal neighborhood constraints.	63
Table 3.3:	Placement results.	65
Table 3.4:	Review of the challenges targeted by physical synthesis algorithms. .	70
Table 4.1:	Summary of the worst case depositing charge for some technology processes for particles of with an LET of $15 MeVcm^2/mg$	75
Table 5.1:	Area and delay of the library cells INV, NAND and NOR in comparison with CWSP cells automatically generated with the tool presented in [LDGR03].	84
Table 5.2:	Total area and delay of CWSP cells.	85
Table 5.3:	Area overhead comparison of TMR and CWSP d-flipflops.	85
Table 5.4:	CWSP and TMR techniques on microprocessors.	87
Table 5.5:	Review of the techniques presented in this chapter.	88
Table 6.1:	Probability of a node as a function of the gate equation [JJ05].	90
Table 6.2:	Approximation of intrinsic MOS gate capacitance.	93
Table 6.3:	The proposed transistor sizing to single event transient attenuation. Results show the area, timing and average power overhead for symmetric and asymmetric sizing techniques for particles with charge $Q = 0.3pC$ with final sensitivity of 50%.	98

Table 6.4:	The proposed transistor sizing to single event transient attenuation. Results show the area, timing and average power overhead for symmetric and asymmetric sizing techniques for particles with charge $Q = 0.3pC$ with final sensitivity of 0%.	98
Table 6.5:	A comparison among TMR, CWSP and the proposed sizing techniques. Results show area overhead and timing penalties to protect some circuits against particles with charge $Q = 0.3pC$ with final sensitivity of 0%.	99

LIST OF ALGORITHMS

1	The proposed <i>superlib</i> generation process.	28
2	The transistor level optimization process used by <i>T-Factor</i>	32
3	The Punch++ layout generation.	39
4	The genetic algorithm.	58
5	An example of layout representation in linear equalities and inequalities. .	66
6	An example of layout representation in linear equalities and inequalities with an objective function.	67
7	A compaction algorithm using linear programming.	68
8	The transistor sizing for SET attenuation.	96
9	Le sizing de transistor pour l'atténuation de SET.	136

ABSTRACT

Deep submicron (DSM) technologies have increased the challenges in circuit designs due to geometry shrinking, power supply reduction, frequency increasing and high logic density. The reliability of integrated circuits is significantly reduced as a consequence of the susceptibility to crosstalk and substrate coupling. In addition, radiation effects are also more significant because particles with low energy, without importance in older technologies, start to be a problem in DSM technologies. All these characteristics emphasize the need for new Electronic Design Automation (EDA) tools. One of the goals of this thesis is to develop EDA tools able to cope with these DSM challenges. This thesis is divided in two major contributions. The first contribution is related to the development of a new methodology able to generate optimized circuits in respect to timing and power consumption. A new design flow is proposed in which the circuit is optimized at transistor level. This methodology allows the optimization of every single transistor according to the capacitances associated to it. Different from the traditional standard cell approach, the layout is generated on demand after a transistor level optimization process. Results show an average 11% delay improvement and more than 30% power saving in comparison with the traditional design flow. The second contribution of this thesis is related with the development of techniques for radiation-hardened circuits. The Code Word State Preserving (CWSP) technique is used to apply timing redundancy into latches and flipflops. This technique presents low area overhead, but timing penalties are totally related with the glitch duration is being attenuated. Further, a new transistor sizing methodology for Single Event Transient (SET) attenuation is proposed. The sizing method is based on an analytic model. The model considers independently pull-up and pull-down blocks. Thus, only transistors directly related to the SET attenuation are sized. Results show smaller area, timing and power consumption overhead in comparison with TMR and CWSP techniques allowing the development of high frequency circuits, with lower area and power overhead.

Keywords: Transistor Level Automatic Layout Generation, Transistor Sizing, Low leakage, Radiation-Hardened.

1 INTRODUCTION

Deep submicron (DSM) technologies have created new challenges to the design of circuits due to geometry shrinking, power supply reduction, frequency increasing and high logic density [CS00]. These characteristics reduce significantly the reliability of the integrated circuits due to the susceptibility to crosstalk and substrate coupling [IEE02].

Interconnections have presented an increased importance in DSM technologies. New technologies have shifted the design paradigm from a conventional logic-dominated to an interconnect-dominated design process [Che06].

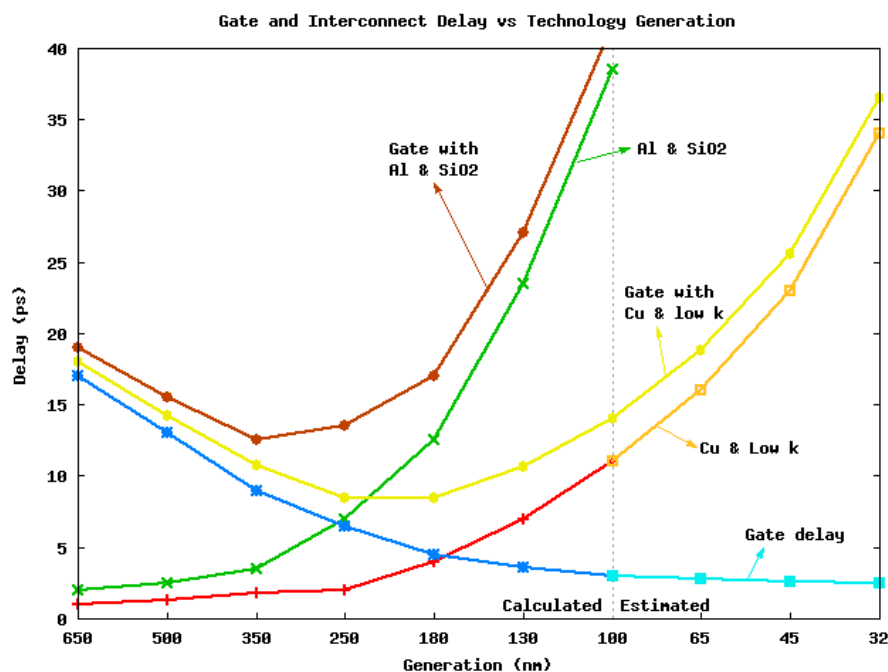


Figure 1.1: Gate and interconnect delay versus technology generation [SIA97]. Delays for feature size below 100nm is estimated.

Figure A.1 shows the amount of timing as a function of the technology generations [SIA97]. Technologies below 100nm are estimated. It is possible to note that the interconnection delay exceeds the gate delay in the 250nm technology process when Al is used for interconnections and SiO₂ is used as dielectric.

Interconnection delay presents more importance in the $0.18\mu m$ technology process when Cu is used in the interconnections and a *low k* dielectric is used as insulator.

The amount of delay of interconnections is not the only important factor in DSM technologies. The high logic density and the increased number of metal layer make interconnection characteristics prediction a very hard task. In a $180nm$ technology, for example, the capacitance per unit length may variate up to 35 times [VWSS04].

Usually, the logic of a circuit is optimized with respect to timing, area and power with assumed capacitances, but their actual values are not known until after the layout phase. Timing closure cannot be achieved with this inaccurate prediction of interconnections without the integration of the layout with the whole process.

The advent of deep submicron technologies also makes mandatory the control of the static power consumption. Designers have been concerned with the dynamic power consumption over all the technology evolution. On the other hand, the power leakage was not taken into account due to its low amount in the total power consumption. However, static power is drastically increasing in DSM technologies.

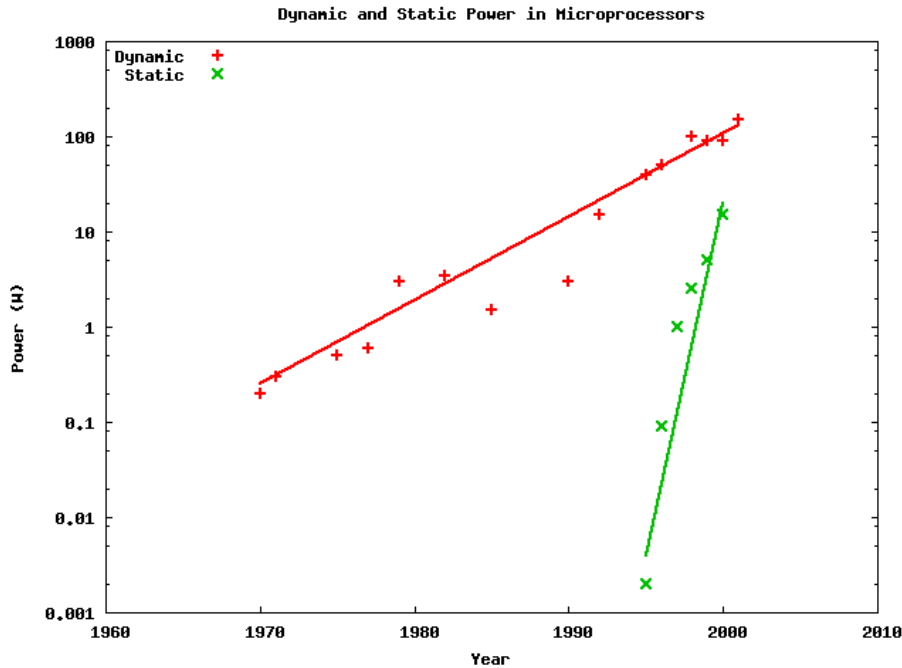


Figure 1.2: Active and static power in microprocessors [Moo03].

Figure A.2 shows the dynamic and static power in microprocessors as a function of the year [Moo03]. The power leakage is projected to exceed the dynamic power consumption in $65nm$ technologies [KAB⁺03].

Dealing with these deep submicron challenges require careful planning [OCG02] and emphasize the need for electronic design automation (EDA) tools able to automatically generate and validate integrated circuits.

Standard cell based designs have dominated the layout generation of digital VLSI circuits due some virtues [GR01]. Standard cells hide the increasingly unpleasant details of

shape-level design rules, IO pins are arranged on individual gates in geometry accessible locations, cells are assembled with relative ease into row-based blocks and cells can be pre-characterized for timing and power.

However, the gap between the standard cell and custom applications are well known in the literature [AND98, ELEHS03].

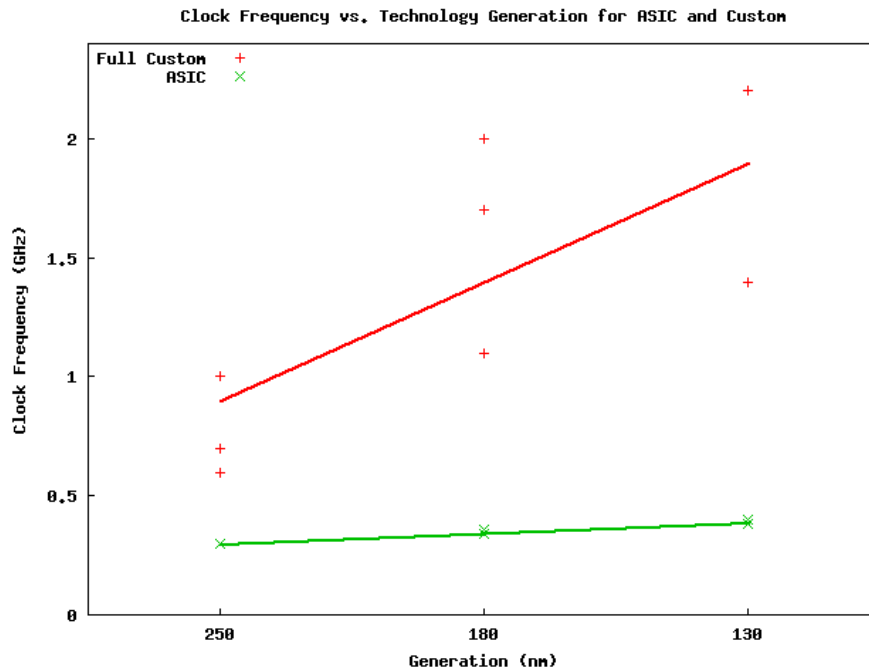


Figure 1.3: Clock frequency of high performance ASIC and custom processors in different process technologies [CK02].

Figure 1.3 quantifies the differences between ASIC and custom chips speeds. Chinnery and Keutzer classify some ASIC and custom processors and evaluates the frequency as function of the generation of technology [CK02]. Some of the custom circuits are Intel, AMD and IBM Power PC processors. Tensilica, Lexra and ARM are classified as ASIC processors.

Figure shows that the frequency gap between custom and ASICs is increasing. Authors report some of the factors that contribute to the superior performance in custom applications. Factors are mainly related to microarchitecture modifications, efficient clock tree design, cells and wire design and transistor sizing. Besides, custom design offers the advantage of full control over the size and the location of each transistor for performance tuning.

Although standard cells are effective to address the layout generation problem, some works have proposed some changes in the traditional design flow in order to cope with the gap between ASIC and custom circuits.

A post-layout optimization is presented in [HO01] in which they propose a downsizing methodology with interconnects preserving. They obtained a 65% power reduction without increasing the circuit delay.

Vujkovic *et al* [VWSS04] show that a standard cell library usually cannot handle the huge connection capacitance variation with a small number of versions of a cell. Furthermore, if a cell is able to deal with the required load capacitance, the cell wastes a lot of power due to the oversized transistors.

The concept of *flex cells* is presented in [RBB05]. They propose the post-layout identification and optimization of a minimum number of critical cells aiming at enhancing the target design performance.

The first contribution of this thesis is to propose a new paradigm in the design of circuits. In the traditional standard design flow, the layout of cells are generated and characterized for timing, area and power consumption. These cells are grouped in a library and used during the whole design process.

We propose a transistor level design flow in which the layout generation is integrated within the whole design process. Transistor level design flow consists on tailoring which individual gate of a circuit. Thus, transistors in critical paths of the circuit are optimized but gates in non-critical paths are maintained with minimum transistors sizes.

The reduction on the dynamic power consumption is a direct consequence of the optimized sizing. Transistor level optimization includes a gate length biasing methodology in order to cope with the power leakage. The length of the transistors outside the critical path are increased to reduce the static power without causing timing penalties in the circuit.

Radiation effects have been also increased in DSM technology process. Reduced transistor gate dimension and low voltage levels cause an increase in the soft error failure rate in integrated circuits.

In the 20th century, single event effects (SEE) were assumed to be concerned mainly in the space environment. Thus, hardening techniques were employed to these applications to avoid loss of information or functional failure.

However, the technology scaling has reduced the reliability of integrated circuits concerning radiation effects. The energy of neutrons, for example, varies between $1MeV$ to $10MeV$ at the atmosphere [Nor01]. Worthless in older technologies, the energy in the atmospheric neutron flux is enough to drastically affect the functionality of current integrated circuits.

Kastensmidt reports that memory elements in a $0.25\mu m$ technology and combinational logic composed of transistors with length smaller than $0.13\mu m$ may be subject to SEE while operating in the atmosphere [Bau05]. For this reason, the development of electronic design automation tools able to cope with the SEE is strongly encouraged.

The second contribution of this thesis is related to the generation of radiation-hardened circuits. A transistor level design flow allows the development of any structure or methodology at transistor level. Different from the conventional standard cell, new techniques could be directly applied in the design flow.

A new technique for protecting sequential elements is proposed in which temporal redundancy is inserted inside the latches and flipflops structures. Besides, a new analytical transistor sizing methodology is proposed in order to cope with single event effects in combinational blocks.

The main characteristic of this new sizing model is the possibility to optimize independently pull-down and pull-up structures of each gate. This allows radiation-hardened combinational circuit to be performed with smaller penalties to the circuit functionality.

1.1 Organization of this Thesis

This thesis is organized in two major parts. The first part is related to the generation of optimized circuits at transistor level.

The proposal of a transistor level design flow is presented in Chapter A.2. Main aspects about the proposed methodology and its impact on the generation of integrated circuits are discussed. A comparison with the traditional design flow is also presented in which some results concerning timing and power consumption are reported.

Algorithms for physical synthesis are presented in Chapter 3. This chapter presents algorithms implemented in the layout generator. Basically, algorithms are used to place, route and compact the layout. The chapter includes a discussion about the importance of the EDA tools on the design of integrated circuits.

The second part of the thesis is related to the protection of integrated circuits against SET. A discussion about radiation effects are given in Chapter 4. The discussion is focused mainly in the Soft Single Event Effects (SEE) due to their relevance with the development of this thesis. The main principles about Single Event Transients (SET) and Single Event Upsets (SEU) are highlighted.

Based on these definitions, Chapter A.3 presents some state-of-the-art techniques for soft error protection. These techniques involve methodologies for radiation hardening in sequential and combinational circuits. A special attention is given to the Code Word State Preserving (CWSP) technique [Ang00] because it was the basis for the development of a new technique aiming at protecting sequential elements.

The last chapter proposes a new transistor sizing methodology for the protection of combinational circuits against SEE. The main contribution of this methodology is the reduced overhead. The main characteristic of the proposed methodology is to find the smallest transistor widths to attenuate SETs in the nodes of a combinational circuit. Another important point is that pull-up and pull-down transistors are independently sized, minimizing the area overhead and power consumption.

2 A TIMING CLOSURE DESIGN FLOW USING A TRANSISTOR LEVEL AUTOMATIC LAYOUT GENERATOR

2.1 Introduction

Traditionally, standard cell libraries are used in digital circuits due some virtues such as the possibility to hide unpleasant details of layout rules and the pre-characterization for timing and power [GR01]. Cells are implemented in different versions allowing to attack different drive strenghts. Usually, a standard cell library presents around ten versions of inverters/buffers and only four or five versions of the other cells [VWSS04].

The advent of deep-submicron technologies shifted the design paradigm from conventional logic-dominated to an interconnect-dominated design process [Che06]. Thus, the huge variance in the wire capacitances cannot be efficiently handled with the limited number of versions of a cell in a standard cell library. In order to drive the output load respecting the required timing, cells may waste a lot of power due to oversized transistors [VWSS04].

In the last years, some works have been presented in order to deal with this layout optimization problems. Vujkovic *et al* presented in [VWSS04] a design flow, which the number of versions of each standard cell is increased. Obtained results show an improvement of up to 20% related to the energy times delay squared (EDD) measurements. The authors consider that the best performance metric is the effectively EDD.

The concept of *flex cells* is presented in [RBB05]. They proposed the identification and optimization of critical cells. In this context, the synthesis of a minimum number of cells is performed aiming at enhancing the target designs performance. For this, the process must take into account both functionality and timing contexts in which each unique cell is used.

A post-layout optimization is presented in [HO01]. They proposed a method for reducing the power consumption by downsizing transistors preserving interconnections. Results presented by authors show that power can be reduced by 65% on average without delay increase when applying post-layout transistor sizing.

The main idea of the proposed work is to explore switch level characteristics in order to reduce the gap between standard cell and full custom design. For this, a complete RTL-to-Layout design flow generation is proposed. The design flow includes the generation of a database used as basis to the logic synthesis and the layout generation. The design flow is based on academic and commercial tools aiming at achieving timing closure and

reduced power consumption.

Timing closure may be obtained by sizing transistors of a circuit in a wide range of possibilities. In other words, a transistor level design flow allows to tailor transistors in critical paths of the circuit maintaining other paths with minimum transistors sizes. The reduction on the power consumption is a direct consequence of this optimized sizing.

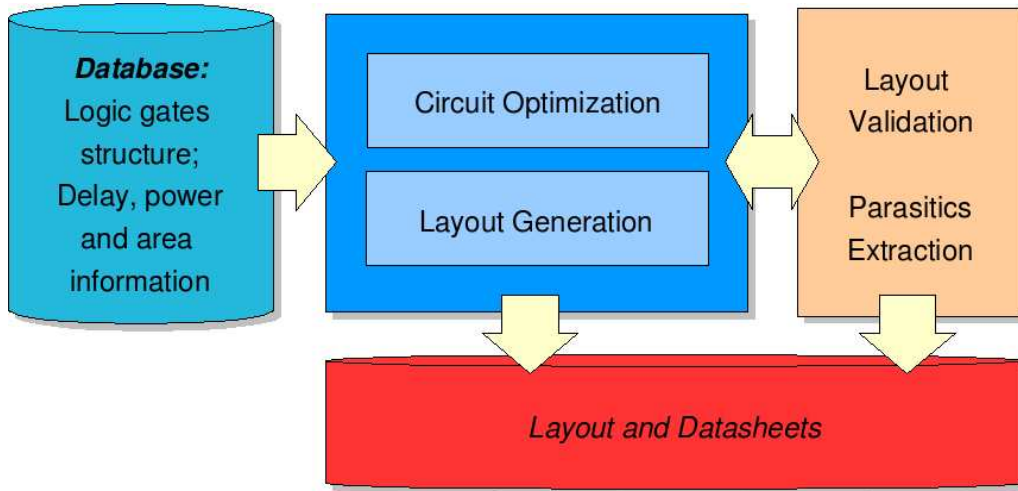


Figure 2.1: The proposed design flow overview.

Figure 2.1 shows an overview of the proposed design flow. First of all, the design flow is based on a database where the structure of each cell, its estimated area, timing and power consumption are stored.

Circuit optimization is the first step in our design flow. Thus, the circuit description is taken with minimum transistor sizes. For each iteration, the longest path is extracted, transistors are resized and the power consumption is analyzed. A power-delay tradeoff can be plot in the end of the optimization process. These preliminarily results allow the designer to choose the best delay and power characteristics according to the design specifications.

The next step is the layout generation. Once the power-delay trade off was chosen by the designer, the layout can be generated. In this step, transistors are placed, folded and routed according to state-of-the-art algorithms. Details about these algorithms are presented in Chapter 3.

Layout extraction allows the evaluation of the generated layout and give to designers important information about the optimization process. If the extracted delay and power met the designer specifications, the flow is finished. Otherwise, extracted data are used to a new optimization phase and the layout is generated again with the optimized layout.

Further details about the design flow are given in the next sections.

2.2 The *Super Library* Generation

Timing closure can be a hard task due to the limited number of logic functions and drive strengths available in standard cell libraries of traditional design flows. Libraries

have up to 200 different logic functions and around four drive strengths. Inverters and buffers usually have up to ten drive strengths. The total number of cells is approximately 3,000 cells. These characteristics are hard limitations to the efficiency of the synthesis and sometimes may degrade power and timing results.

In a standard cell design flow, cells are generated and characterized for timing and power. The whole process takes long time and needs an enormous effort from designers. As consequence, a limited number of cells and drive strengths are available for logic synthesis and technology mapping.

Table 2.1: Maximum number of logic functions obtained by stacked transistors [DRSVW87].

		Number of PMOS stacked transistors				
		1	2	3	4	5
Number of NMOS stacked transistors	1	1	2	3	4	5
	2	2	7	18	42	90
	3	3	18	87	396	1677
	4	4	42	396	3503	28435
	5	5	90	1677	28435	425803

A research published in [DRSVW87] shows the possibilities on generating libraries according the number of stacked transistors. Table 2.1 shows the result of this research where the number of possible combinations are very larger than the number of logic functions available in a library.

Another aspect concerning the number of library cells is related to the quality of the algorithms for logic and physical synthesis. Algorithms presented in the literature in the last years usually show a linear behavior in respect to the complexity of circuits. This linear behavior allows to increase significantly the number of cells in the libraries.

Figure 2.2 shows the performance and occupied memory of synthesis tools as a function of the number of cells in a library. Studies of the globally based synthesis algorithms show that they can effectively take advantage of sets of 10K or more cells due to the linear runtime and memory usage [CAD07a].

The concept of *library-free mapping* was introduced in [RRAR97]. The library-free is a method for mapping a set of boolean equations into a set of static CMOS complex gates (SCCG) under a constraint in the number of stacked transistors. Reis highlights that the number of transistors in synthesized circuits is inversely proportional with the static power consumption.

The most important characteristic of the library-free mapping is related to the reduced number of transistors in comparison with standard cell libraries with a restricted number of logic functions. The library free mapping may result in 20-30% reduction in the number of transistors.

On the other hand, a library-free methodology may not lead to good results due to the absence of timing and/or power information. It is clear that the synthesis cannot be efficient with this lack of information. Although the library-free mapping presents good characteristics concerning the reduced number of transistors, timing closure cannot be

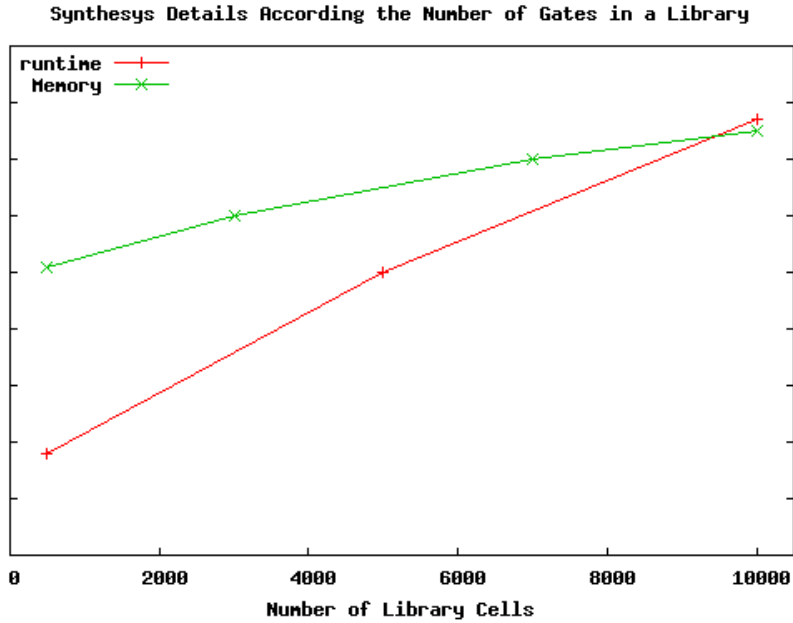


Figure 2.2: Runtime and occupied memory as a function of the number of gates in a library [CAD07a].

reached without the delay information of each cell.

For this reason, the development of a *super library* is proposed in this thesis, which consists on developing a library with a very large number of cells enriched by the timing information. Thus, the synthesis is able to generate a circuit with small number of transistors as well as targets timing closure.

The cells layout does not exist in the *superlib*. Only the logic function, cells structure (transistors and connections) and the estimated timing and occupied area are known. These informations are enough to allow timing closure logic synthesis.

The ability of synthesis algorithms on exploring a wide number of logic functions is not so clear in the literature. One important question about the efficacy of our *superlib* was whether commercial tools really take advantage of a wide number of logic functions at the logic synthesis.

Figure A.3 shows that commercial tools are able to explore the synthesis when a wide number of logic functions are available. These results were obtained with Cadence RTL Compiler [CAD07b]. Analyzing these ten circuits, we note that circuits mapped with our *superlib* have around 52% more complex logic functions. Simple gates such as NANDs, NORs, inverters and buffers are not included in these results.

2.2.1 The Development of a *superlib*

The *superlib* is generated in Synopsys Liberty library format [SYN07] and contains delay information about different versions of each logic function. This library is composed of 3,503 different logic functions, which is composed by every logic function with up to four stacked transistors, as shown in Table 2.1.

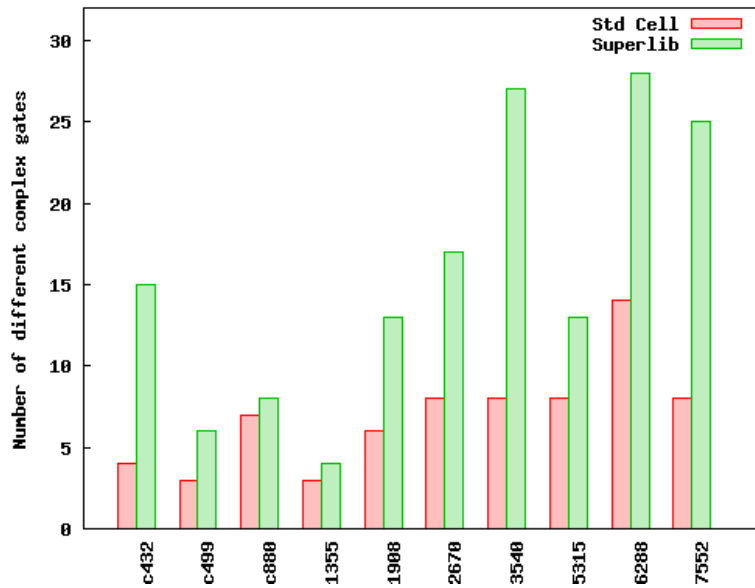


Figure 2.3: Number of different complex logic functions in circuits in a commercial $0.35\mu m$ technology. Results were obtained with Cadence RTL Compiler [CAD07b].

Each logic function was implemented in four different drive strengths by sizing its transistors from 1 to 4 times the minimum possible transistor width. Inverters and buffers were implemented with 8 and 32 different versions, respectively. The resulting library is composed by almost 15,000 cells.

The number of drive strengths of the *superlib* was limited due to the runtime and memory usage. Although the linear behavior of new synthesis tools, very large libraries may take long time to converge to a solution. In other words, the number of drive strengths was limited in order to deal with the tradeoff between quality of results and runtime during the synthesis.

Table 2.2: Synthesis runtime of the ISCAS85 benchmarks with a commercial $0.35\mu m$ standard cell library and the *superlib* with Cadence RTL Compiler [CAD07b].

Benchmark	Standard Cell	superlib
c432	3s	2min21s
c499	9s	2min31s
c880	8s	1min49s
c1355	9s	2min28s
c1908	8s	1min34s
c2670	14s	2min12
c3540	22s	4min38s
c5315	32s	7min52s
c6288	47s	26min30s
c7552	48s	14min36s

Table 2.2 shows a comparison between the runtime of the logic synthesis with a standard cell library and the *superlib* for the ISCAS85 benchmarks. We consider a standard cell library with around 500 cells while the *superlib* is composed by 14,048 cells in this comparison. The drawback of our methodology is visible when runtime is considered and these results justify the limitations we impose to the *superlib* concerning the number of available cells.

This limitation is attenuated by sizing the transistors just after the synthesis. Further details are given in Section A.2.2.

One important remark is that we do not take into account the power consumption at this stage. This is done in order to reduce the library generation execution time. The power consumption is proportional to the transistors area. Thus, this information may be omitted in the library characterization.

Algorithm 1 The proposed *superlib* generation process.

Require: Set of cells C , Set of input slopes S , Set of output Capacitances O

Ensure: Liberty file

```

1: for all  $c \in C$  do
2:    $T \leftarrow \text{getTransistors}(c)$ ;
3:    $I \leftarrow \text{getInputs}(c)$ ;
4:    $G \leftarrow \text{createGraph}(T)$ ;
5:   for all  $i \in I$  do
6:     {Finding the rise time from the input  $i$  to the cell output}
7:      $D_{rise} \leftarrow \emptyset$  { $D_{rise}$  are rise times to each input slope and output capacitance}
8:      $P \leftarrow \text{findTransistorsInPullUpPath}(T, i)$ ;
9:      $N \leftarrow T \setminus P$ ;
10:     $\text{connectInputsToVDD}(N)$ ;
11:     $\text{setInputSwitchingFromOneToZero}(P)$ ;
12:    for all  $s \in S, o \in O$  do
13:       $d \leftarrow \text{getDelay}(G, P, N, s, o)$ ; {Run hspice}
14:       $D_{rise} \leftarrow D_{rise} \cup \{d\}$ ;
15:    end for
16:    {Finding the fall time from the input  $i$  to the cell output}
17:     $D_{fall} \leftarrow \emptyset$  { $D_{fall}$  are fall times to each input slope and output capacitance}
18:     $P \leftarrow \text{findTransistorsInPullDownPath}(T, i)$ ;
19:     $N \leftarrow T \setminus P$ ; { $N$  is the set of transistors outside the biggest path}
20:     $\text{connectInputsToGND}(N)$ ;
21:     $\text{setInputSwitchingFromZeroToOne}(P)$ ;
22:    for all  $s \in S, o \in O$  do
23:       $d \leftarrow \text{getDelay}(G, P, N, s, o)$ ; {Run hspice}
24:       $D_{fall} \leftarrow D_{fall} \cup \{d\}$ ;
25:    end for
26:  end for
27: end for

```

Algorithm 1 presents the method for generating the *superlib*. Considering a library composed by a set of cells C associated with a set of input slopes S and a set of output

capacitances O , the algorithm simulates the spice netlist of each cell for rise and fall time. The timing characterization of the cells is done by electric simulations with Synopsys HSPICE [SYN07].

In order to reduce the number of simulation runs in the characterization process, we simulate only input vectors that stimulate the biggest path between supply source and the cell output for each cell input. We describe a logic function as a graph G (line 4), where each variable of the logic function is represented by an input signal (vertex in the graph). Drain/Source nodes are represented by edges.

The algorithm searches the biggest path crossing the variable i between the VDD/GND and the output. After the biggest path is assumed to be known (line 8 for rise time, and line 18 for fall time), transistors in the path are configured to switch from the “OFF” state to the “ON” state (line 10 and 20) and transistors outside the biggest path are set to the “OFF” state. Thus, for each input signal of a cell c , only one input vector is necessary for measuring the rise time and another input vector for measuring the fall time.

Once the path is known and input signals are defined, a spice simulation is done for each input slope $s \in S$ and each output capacitance $o \in O$ (lines 12-15 and 22-25). At the end, the sets of delay times D_{rise} and D_{fall} contain the whole information about the cells timing.

Figure 2.4 shows the process of generating the propagation time of each cell. Only the pull-up is represented in the figure. Considering the transistor schematic in Figure 2.4(a), the graph is created. The graph is used to find the longest path between the V_{DD} and the cell output (Figure 2.4(b)). An example of the representation of the rise and fall time in liberty format is shown in Figure 2.4(c). The field `values` is a 2D table and contains the delay information as a function of input slopes (rows) and output capacitances (columns).

The whole process (more than 750,000 simulations) was automated by scripts. The generation of the *superlib* took around 6 days in a SunfireV890 with 8Gb of RAM.

2.3 The Transistor Level Design Flow

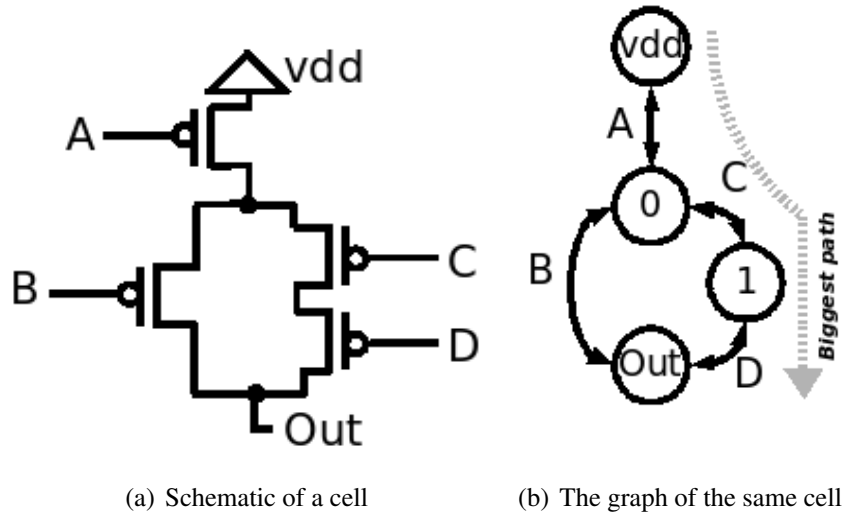
The transistor level optimization has been considered in several works when timing closure is achieved [HO01, VWSS04, RBB05]. Vujkovic *et al.* report in [VWSS04] that the capacitance per unit length varies around 35 times in a $0.18\mu m$ technology. This shows how the optimization of each cell as a function of the output load capacitance is important, specially in the critical path.

In this section we present a complete design flow based on transistors optimization. The main goal of this methodology is to explore the capabilities of a transistor level design flow to deal with the challenges present in deep submicrom technologies.

Once the *superlib* was created, the layout of the circuit can be generated. Figure A.4 shows in details the design flow proposed in this work. The proposed transistor level design methodology consists on a set of commercial and academic tools that deal with the timing closure challenges at the same time as it offers to designers a complete RTL-to-GDS design flow. Developed tools are labeled with letter D.

Main differences of the traditional standard cell flow are the following:

- The possibility to synthesize a circuit with a wide range of cells by using the *super-*



```

timing() {
  related_pin    : "A";
  timing_sense   : positive_unate;
  cell_rise(del_0_4_5) {
    values( " 0.177, 0.206, 0.293, 0.614, 1.036",\
            " 0.221, 0.247, 0.338, 0.657, 1.079",\
            " 0.240, 0.268, 0.357, 0.678, 1.099",\
            " 0.233, 0.264, 0.355, 0.675, 1.097");
  }
}

```

(c) Example of a cell representation in liberty format

Figure 2.4: Generation of the liberty file. This file is used later in the logic synthesis.

lib (Figure A.4 Label A);

- A transistor level optimization tool called **T-Factor** able to size each transistor of the circuit in order to deal with timing closure (Figure A.4 Label B);
- A new transistor level layout generation tool called **Punch++** able to generate any kind of Static CMOS gate (Figure A.4 Label C).

The first step in the design flow is to generate a spice-like netlist of the circuit. In the proposed design flow, circuits described in high level languages such as VHDL or Verilog are converted to netlist descriptions by using Cadence RTL Compiler [CAD07b]. After synthesis, the verilog netlist is converted to a spice netlist description.

The transistor level optimization starts with the spice netlist description. The tool **T-Factor** finds transistors in the critical paths and these transistors are sized in order to meet timing specifications while transistors in less important paths are maintained with minimum sizes.

The cells placement is based on area estimates with the Cadence Amoeba placer [CAD07b]. After the placement, the layout of the entire circuit is generated and routed by

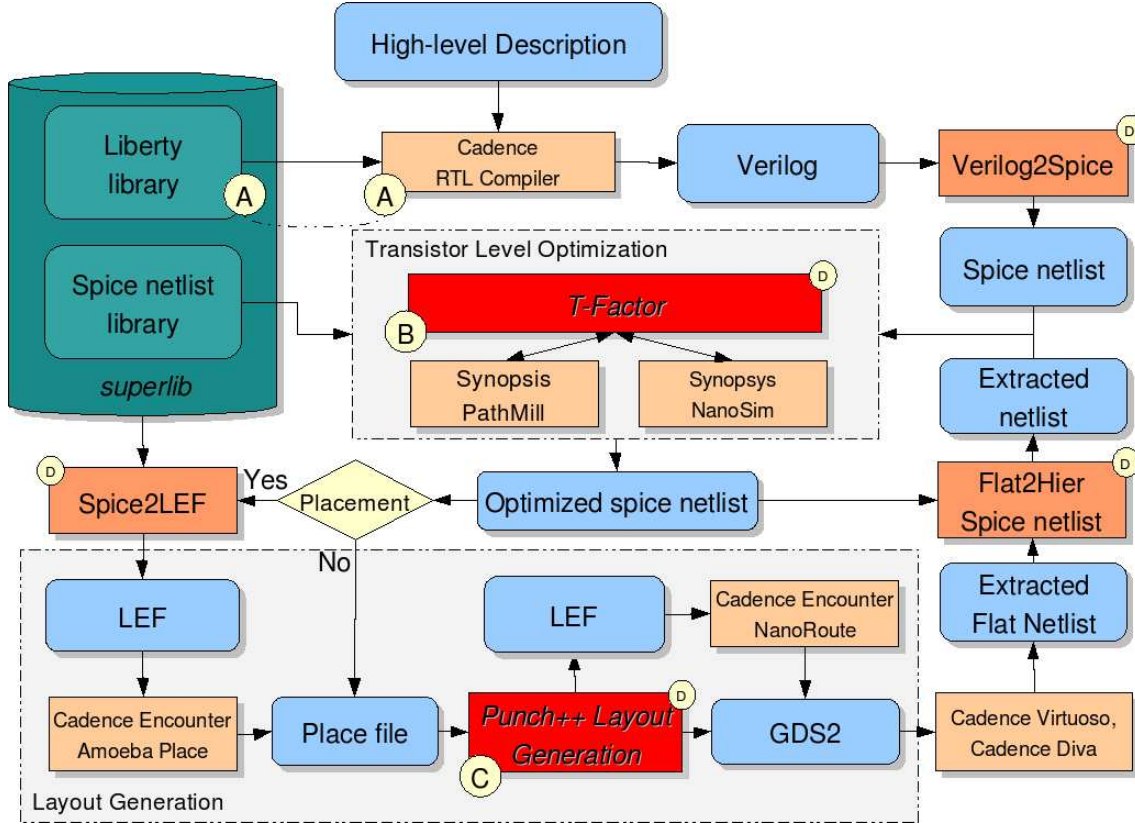


Figure 2.5: The proposed transistor level design flow. A) Logic synthesis with the *superlib*. B) Transistor level optimization tool. C) Layout generation tool, D) Developed tools to cope with the proposed transistor level design flow.

the Cadence Nanoroute. The last step of the design flow is the DRC, LVS and parasitics extraction, which makes possible to evaluate the correctness of the circuit.

Details about the most important steps in the transistor level design flow are discussed in the following.

2.3.1 The Transistor Level Optimization

As previously discussed, the number of cells was limited to 3,503 targeting runtime reduction in the logic synthesis. The number of drive strengths was also limited by the same reason.

The limitation concerning the number of drive strengths is compensated by optimizing the circuit with a transistor sizing tool. *T-Factor* works in association with the Nanosim e Pathmill from Synopsys [SYN07] in order to optimize the circuit at transistor level. Dynamic power simulation is performed by Nanosim and static timing analysis is done by Pathmill.

The transistor level optimization is shown in Algorithm 2. The optimization process starts with the generation of a set of input vectors V (line 3). These vectors are used to analyze the dynamic power consumption. The number of input vectors are defined by s because the power analysis runtime is directly proportional to s . A smaller s may be defined to big circuits in order to reduce the runtime when analyzing the power consump-

Algorithm 2 The transistor level optimization process used by *T-Factor*.

Require: The netlist spice N , Number of iterations k , Maximum transistor width m , Transistor size step p , Input vector size s

Ensure: Set of optimized netlists N_{new} with timing and power information

```

1:  $N_{new} \leftarrow N$ ;
2:  $i \leftarrow 0$ ;
3:  $V \leftarrow \text{generateInputVectors}(s)$ ;
4: while  $i < k$  do
5:    $pm \leftarrow \text{getCriticalPath}(N)$ ;           {run Pathmill}
6:    $ns \leftarrow \text{getPowerInformation}(N, V)$ ; {run Nanosim}
7:    $c \leftarrow \text{findCellWithWorstLoadFactor}(pm)$ ;
8:    $T \leftarrow \text{findTransistorToSize}(c)$ ;
9:   for all  $t \in T$  do
10:     $w \leftarrow t.w + p$ ;
11:    if  $w > m$  then
12:      stop; {Stop and go to line 21}
13:    else
14:       $t.w \leftarrow w$ ;
15:    end if
16:  end for
17:   $c_{new} \leftarrow \text{updateTransistors}(c, T)$ ;
18:   $N_{new} \leftarrow N_{new} \cup \{c_{new}\} \setminus \{c\}$ ;
19:   $i++$ ;
20: end while
21: plotDelayPowerTradeoff();

```

tion.

Function `getCriticalPath(  $N$  )` (line 5) consists on running Pathmill and extracting the critical path. Pathmill reports the worst path with the delay in each stage and the capacitances of each node. The power consumption is detected by the function `getPowerInformation(  $N, V$  )` (line 6) as a function of the input vector V .

The load factor is used to find a cell among all the candidates to sizing. Load factor is used in many works as criterion to size cells in circuits [CHP00, SWL⁺03]. The load factor is defined by

$$F_{load} = \frac{C_{load}}{C_{in}} \quad (2.1)$$

where C_{load} is the load capacitance on the cell outputs and C_{in} is the input gate capacitance. The candidate cell is given by the function `findCellWithWorstLoadFactor(  $pm$  )` at line 7 where the cell with the worst load factor is chosen to be sized.

Function `findTransistorToSize(  $c$  )` finds the transistors path between the supply line to the cell output by a graph structure similar to Figure 2.4(b). Lines 9 to 16 increase the width of every transistor $t \in T$ by s .

After the transistors sizing, the netlist is updated and the whole process is repeated until the number of iterations meet the maximum number of iterations k defined by the

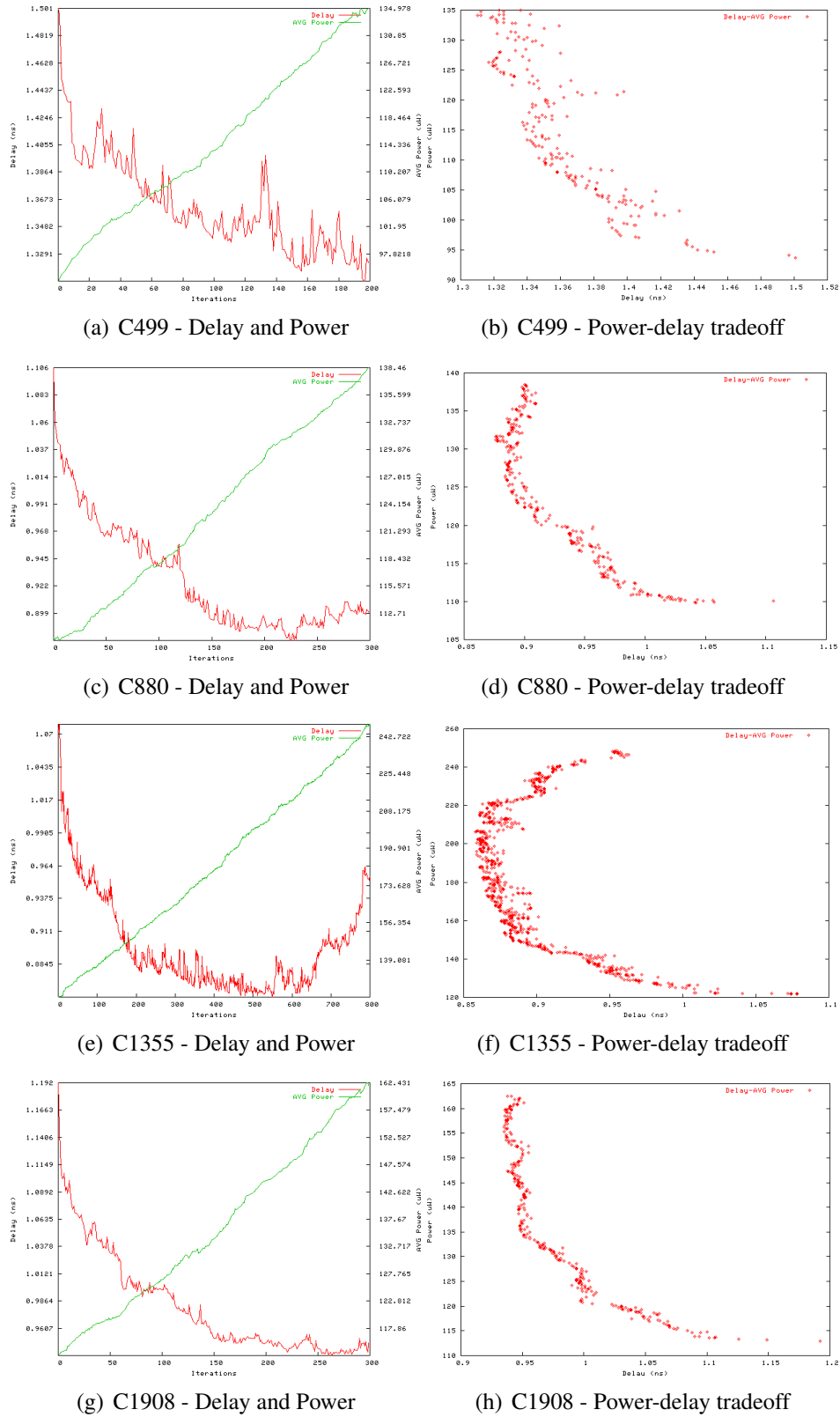


Figure 2.6: Delay and power consumption results for some circuits from the ISCAS'85 benchmarks.

designer or the gates have the maximum width m .

Figure A.5 shows the delay and power consumption as result of the first optimization step for some ISCAS'85 circuits. Figures (a)(c)(e) and (g) present the delay and power consumption for each iteration. One important aspect is the linear grown of the power consumption as a function of the sizing algorithm.

Huge amount of power is inevitable when searching fastest circuits. In most designs, the smallest delay is not the best option due to the increased power consumption and circuit surface. Furthermore, aiming at low power devices, the consequence is a bigger delay of the circuit.

Sometimes the criterion for designing a circuit is neither an extremely low power nor a high performance, but a good tradeoff between delay and power consumption. Thus, Figures A.5 (b)(d)(f) and (h) show the power-delay tradeoff for the circuits where the designer is able to evaluate the optimization process. Each point in the power-delay plot corresponds to a possible circuit design.

The efficiency of the transistor optimization process cannot be measured while analyzing only the worst delay of a circuit. The efficient is proven by analyzing all paths of the circuit.

The operation frequency of a combinational circuit is given by the worst path delay. A big number of paths with very small delay may signify oversized transistors due to an inefficient sizing algorithm. As a consequence, an inefficient sizing strategy may lead to an excessive and undesirable power consumption.

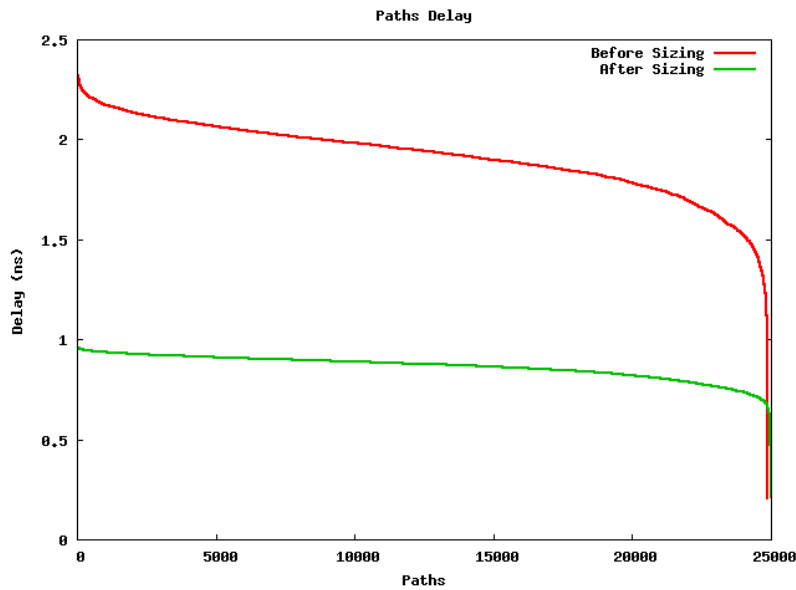


Figure 2.7: Delay paths of the circuit c1355.

Figure A.6 illustrates delay paths before the transistor sizing and after the transistor sizing. The X axis shows all the paths in the circuit c1355 while the Y axis represents the delay of each path. The almost horizontal line given by the delay paths after sizing shows the efficiency of the sizing algorithm. This uniformity in delay paths means that gates in the fastest paths present adequate size and are not over consuming.

2.3.2 Transistor Level Optimization for Leakage Reduction

The static power has become an important amount of the total power consumption. The power dissipation from chip leakage is approaching the dynamic power, and off-state subthreshold leakage is projected to exceed the dynamic power consumption as technology drops below the 65nm feature size [KAB⁺03].

For these reasons, the static power consumption has to be incorporated in the design of systems in the technology process.

Many techniques have been presented in the last years aiming at reducing the static power consumption. The leakage problem has been addressed at design stage by various techniques such as transistor stacking, sleep transistor insertion, V_{DD} assignment and transistor length biasing.

Transistor stacking is the technique to duplicate transistors aiming at reducing the subthreshold leakage when transistors are in standby mode [NDB⁺02]. Figure 2.8 shows an example of stacked transistors structure.

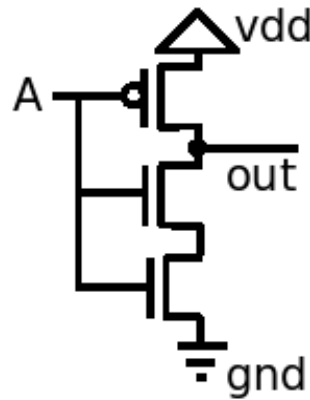


Figure 2.8: An NMOS two-stack inverter reduces the subthreshold leakage when input stays at logic “0”.

Sometimes the stacked transistor widths are divided by two in order to ensure the same input load. Thus, the previous gate delay and the switching power remains unchanged. However, the gate delay is increased and timing closure may not be achieved.

The main idea in the insertion of *sleep transistors* is to cut the subthreshold current when CMOS gates are in standby mode [LH03, BBMM04]. The methodology seems to be very interesting because a whole block can be turned off, but the main drawback of this technique is related to delay penalties.

The activation time needed by the sleep transistors when they change from the state “OFF” to the state “ON”, may be longer than the response time of many cells in the design (specially those placed close to the primary inputs (PI)). This includes an extra time on the circuit.

Figure 2.9 illustrates an example of fine-grained sleep transistor. Usually sleep transistors do not only control the shutoff of a gate but also a set of gates in order to reduce the area overhead.

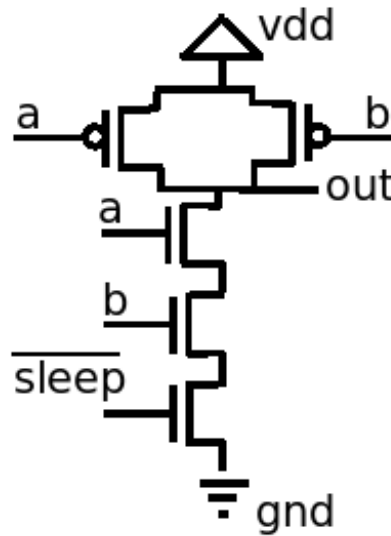


Figure 2.9: An example of fine-grained sleep transistors.

V_{DD} assignment consists of assigning different supply power voltages to different cells in the circuit [LH05]. Thus, cells with lower V_{DD} voltage present a smaller power consumption.

The V_{DD} assignment technique demands careful analysis in relation to the cells placement and supply lines routing. Cells must be grouped in order to reduce the supply routing, increasing the wire length and routing congestion.

Considering that the power planning is already a hard task in submicron technologies due to local hot spots, insufficient power supply and signal integrity problems, V_{DD} assignment insert more challenges in the traditional design flow.

Gate length biasing consists of adjusting the length of the transistors to reduce the power leakage [GKSS04, KMS05, BCV06]. The leakage is inversely proportional with the gate length. However, the delay of a transistor increases with the gate length. Thus, fastest paths can be sized while the transistors in the critical path may maintain its length in order to cope with timing closure.

From the point of view of the transistor level layout generation, all the techniques above can be applied. However, the gate length is a technique with the whole process can be performed at switch level. This technique is very hard to be implemented in the traditional standard cell approach because new versions of each cell must be generated and characterized.

The second possibility to gate length biasing in standard cells is the post layout optimization. The length of the transistor in the fastest paths may be manually modified. Identify and modify the layout is a very complex task and add many hours to design.

For the transistor level layout generation, the gate length biasing is simple to be included in the design flow because the layout is the last step of the whole process. For this reason, we include the gate length biasing in our design flow. Transistors are sized for power leakage reduction after the circuit is sized for timing closure.

Figure A.7 shows the normalized leakage current in deep submicron technologies as

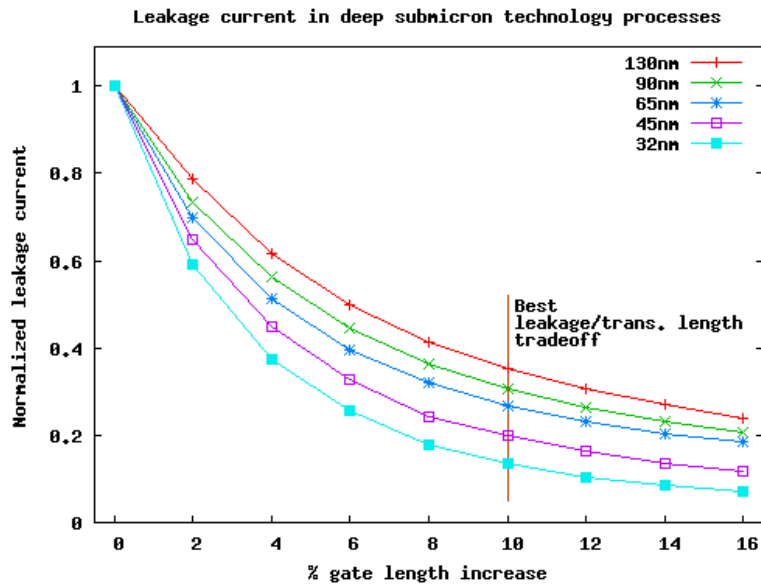


Figure 2.10: Leakage current in deep submicron technology process.

a function of the transistor length biasing. These data were obtained by spice simulations with the predictive technology models presented in [ZC07].

Results show that the transistor length biasing is more efficient with the technology shrinking. However, an upper bound to the transistor length sizing may be defined due to the exponential reduction in the leakage current.

The leakage reduction gain is very small when the transistor length is bigger than 10% for all technology processes. For this reason, we define a 10% upper bound to transistor length sizing for our experiments.

Table 2.3: The gate length biasing technique for some ISCAS'85 benchmarks with a 65nm technology process.

Benchmark	Sized cells		Power leakage		
	#Cell	(%)	Before	After	After (%)
C432	139/209	66%	3.5 μW	2.1 μW	60%
C499	208/296	70%	12.5 μW	8.1 μW	64%
C880	290/359	80%	10.6 μW	5.9 μW	55%
C1355	247/446	55%	10.3 μW	6.6 μW	64%
C1908	245/372	65%	14.4 μW	11.0 μW	76%
C3540	526/704	74%	20.1 μW	12.3 μW	61%
Average leakage after gate length					63%

Results about the gate length biasing are shown in Table A.1. Results show that the gate length biasing is very efficient on the reduction of the power leakage. An average of 70% of the cells were sized and the resulting circuit spend 60% of the initial power

leakage in average. It is important to remark that the delay and power dynamic is not increased.

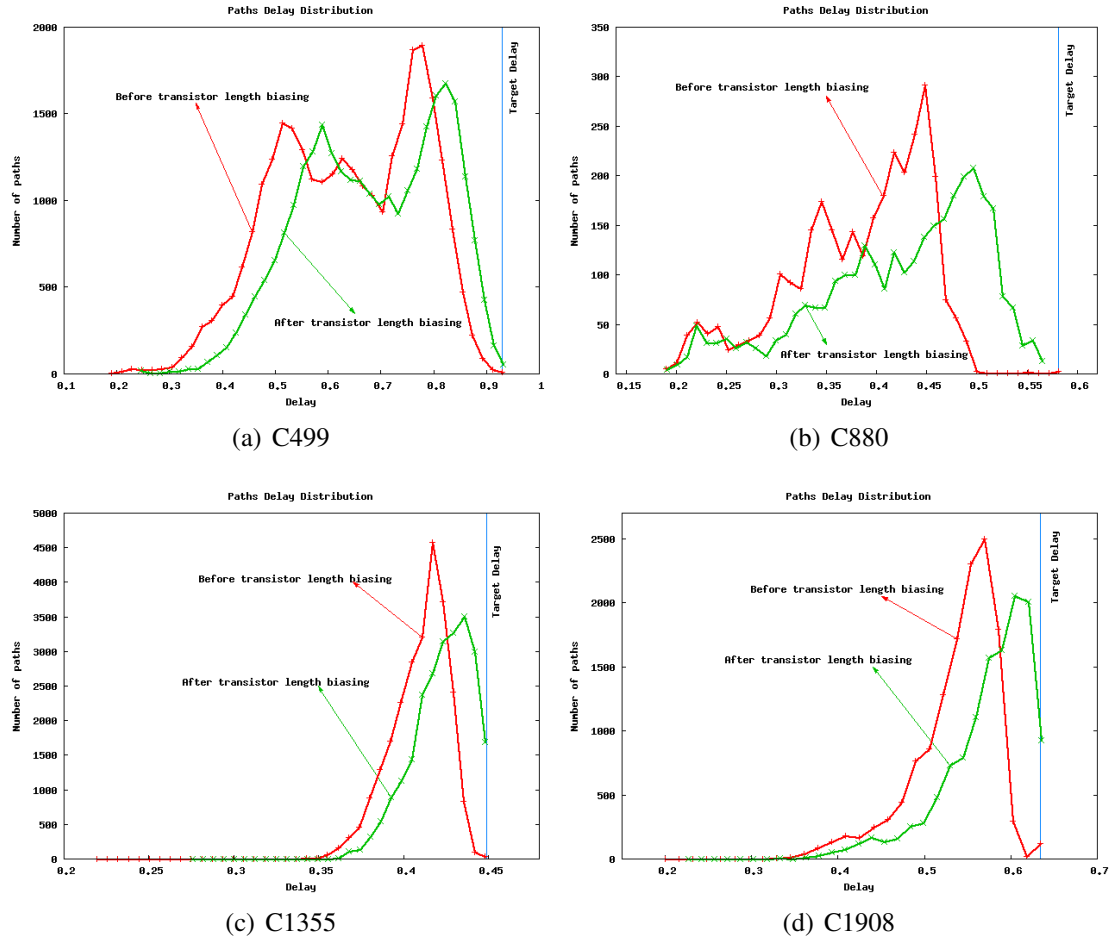


Figure 2.11: Delay paths distribution before and after transistors length biasing to the benchmarks ISCAS'85.

The distribution of the delay paths before and after the transistors length biasing is shown in Figure A.8. Curves show the bigger number of paths close to the target delay after the gate length biasing.

One important remark is related to the process variability and the big number of paths close to the target delay. The variability of the process technology may increase the delay of some paths causing a timing violation. For this reason, worst case corners must be used at simulation time to reduce the impact of the variability in the circuit delay.

2.3.3 The Transistor Level Layout Generation

As previously discussed in this chapter, the layout of the whole circuit is fully generated on demand, without any previously defined layout. Only the netlist with the sized transistors and the technology rules are needed.

The transistor level layout generation involves several tasks as illustrated in Figure A.4. First, the layout cannot be generated without the placement of cells. Thus, estimated

information (delay, power and area) of each cell is used to place the cells in the circuit. In order to estimate this information, the spice netlist of each cell is used. The transistors structure allows the power and timing estimation. The area can be estimated with the cells structure and some technology rules. Cadence Amoeba [CAD07b] is used to place the cells.

After the cells placement, the layout generation can be performed. Algorithm 3 shows the process to generate the circuit layout.

Algorithm 3 The Punch++ layout generation.

Require: Set of cells N in spice format, Technology Rules T , Placement P

Ensure: Layout in GDSII format L

```

1: foldTransistors(  $N$  );
2:  $R \leftarrow \text{applyPlacement}( N, P );$            {  $R$  are the rows }
3: for all  $r \in R$  do
4:    $L_r \leftarrow \text{placeTransistors}( N, r );$ 
5:    $L_r \leftarrow \text{routeTransistors}( L_r );$ 
6:    $L_r \leftarrow \text{compactLayout}( L_r, T );$ 
7: end for
8:  $L \leftarrow \text{routeCircuit}( L );$            { Call Cadence nanoroute }
9: writeLayout(  $L$  );
```

The algorithm starts with the transistors folding. The folding technique consists on breaking big transistors in equivalent parallel transistors (line 1). Usually, big transistors are not easy to place and may increase the area occupied by the cells. For these reasons, the folding technique is applied to the netlist.

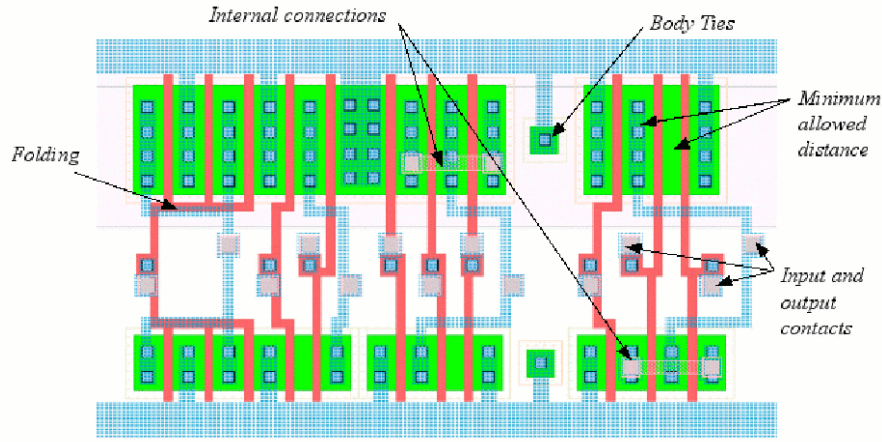
The placement is applied to the netlist in Function `applyPlacement(  $N, P$  )` (line 2) where cells are organized in the set of rows R and placed side by side.

For each row $r \in R$, transistors are placed and routed in Functions `placeTransistors(  $N, r$  )` and `routeTransistors(  $L_r$  )`, respectively. Transistors are placed and routed without taking into account the technology rules. The Function `compactLayout(  $L_r, T$  )` is responsible for compacting the layout and applying the technology rules to the layout. More details about the algorithms used in these three steps are given in Chapter 3.

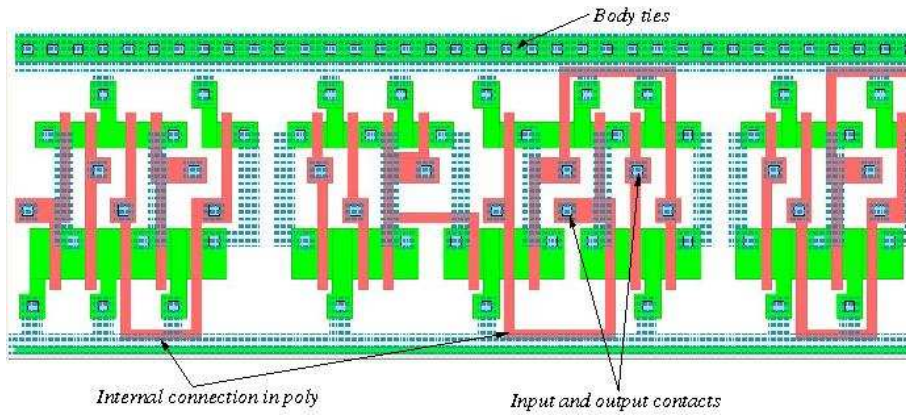
Once the layout of each row is performed, the circuit is routed by the Cadence nanoroute [CAD07b] (Function `routeCircuit(  $L$  )` in line 8) and the circuit is stored in GDSII format (Function `writeLayout(  $L$  )` in line 9). The GDSII is the standard format used by the industry.

When the layout is fully generated, the designer is able to extract the parasitics using Cadence DIVA [CAD07b]. Layout validation is also done with LVS and DRC tools. The extracted netlist is used to evaluate electric characteristics such as power and timing. In addition, if the circuit does not meet the specifications, the designer can repeat the optimization process (Section 2.3.1) in order to improve the electrical characteristics of the circuit.

Figure 2.12 shows a comparison between the layout style of the tool presented in [Laz03] and the layout style of the tool developed in this work. Basically, differences can



(a) Layout from the work proposed in [Laz03]



(b) The new layout style

Figure 2.12: Layout style of two transistor level automatic layout generators. The *Punch* and *Punch++* tools.

be seen in the placement of body ties and interconnections.

The layout style from [Laz03] presents internal connections implemented with the first and second metal layers. The new tool only connects transistor with the first metal layer, but small connection can be performed with polysilicon lines. Details about the algorithms implemented in the new tool are given in Chapter 3.

2.4 Transistor-Level vs Traditional Design Flow Summary

Table 2.4 shows a summary about some characteristics of the traditional standard cell and the proposed transistor level design flow. The main steps in the whole process are discussed. The design flow includes logic and physical synthesis.

2.4.1 A Comparison with the Traditional Standard Cell Design Flow

Table A.2 presents results of the proposed method in comparison with the standard cell approach for a commercial $0.35\ \mu\text{m}$ technology. Results are very interesting because

Table 2.4: Summary of differences of between a traditional standard cell and the proposed transistor level design flow.

Library Cells Generation	Traditional design Flow: The library contains the layout of cells with area, timing and power information. Library usually contains a few versions of each logic function. Proposed design Flow: Layout of the logic functions does not exist. Information concerning timing and power is done based on the spice simulations. A library contains thousands of logic functions.
Logic Synthesis	Traditional design Flow: Synthesis based on the library of cells. Proposed design Flow: Synthesis based on the library of cells.
Transistor Level Optimization	Traditional design Flow: N/A Proposed design Flow: Consists on optimizing transistors for timing and power. Static timing analysis is done to find and optimize critical paths with a bigger range of possibilities. Transistors outside the critical paths are maintained with minimum sizes to reduce power consumption. Power leakage is attenuated by gate length biasing.
Placement	Traditional design Flow: Based on the cells information. Proposed design Flow: Based on the estimated area of cells and timing from electrical simulation.
Layout Generation	Traditional design Flow: N/A. Layout is done with library generation. Proposed design Flow: Layout is generated by algorithms applied to an optimized transistor level netlist.
Routing	Traditional design Flow: Traditional Routing. Proposed design Flow: Traditional Routing.

they show the efficiency of our transistor level design flow.

The design process was done based on high effort to meet the minimum possible delay. We noted that this high effort resulted in the insertion of many buffers in the critical path. This explains the low gain concerning the timing (around 11%) of our methodology in comparison with the standard cell approach. The power consumption gain in these

Table 2.5: Comparison between layouts generated by the standard cell approach and the proposed transistor level design flow.

Circuit	Timing (ns)			Total Power Consumption (uW)		
	Std Cells	Proposed	Gain	Std Cells	Proposed	Gain
C432	3.97	3.68	7.3%	4416	3726	15.6%
C499	2.36	1.89	19.9%	11881	7122	40.0%
C880	1.88	1.85	1.5%	5592	3984	28.7%
C1355	2.50	2.45	2%	12071	6965	42.2%
C1908	2.39	2.06	13.8%	9493	6007	36.7%
C3540	5.15	4.05	21.4%	21141	15235	27.9%
C6288	9.46	7.98	15.6%	211593	145660	31.1%
Average Gain			11.6%			31.7%

circuits is between 15% and 42% due to the number of complex gates in the circuit and the optimized transistor width.

These results were obtained by attempt for improving the layout quality concerning polysilicon and metal connections and by the possibility to optimize the circuit in relation to the wide number of logic functions and drive strengths present in a library.

Figure 2.13 shows two layouts of circuits generated by the proposed transistor level design flow.

2.5 Conclusion

This chapter explores the possibilities to develop a transistor level design flow as alternative to the traditional standard cell. In a traditional design flow, the layout of cells is generated, characterized and grouped in libraries. These libraries contain information about occupied area, timing and power consumption.

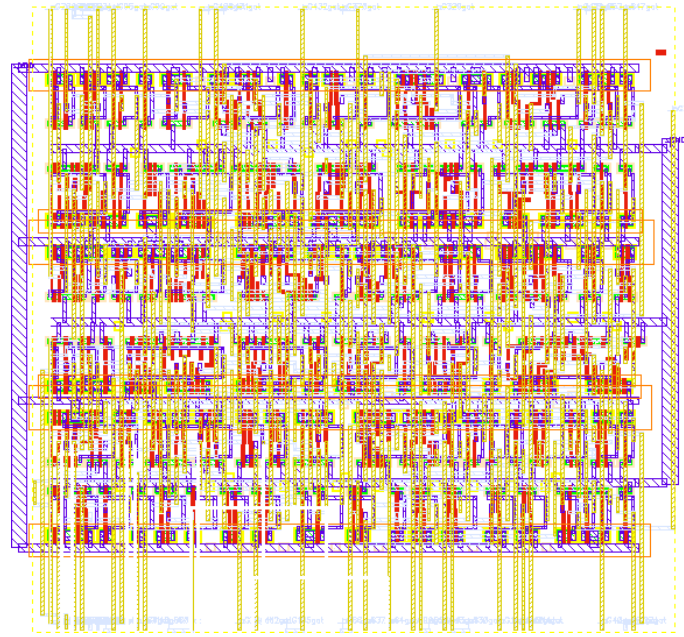
The transistor level design flow is a different paradigm. Libraries do not contain the layout of the cells, but only timing and power information. This information is obtained by spice simulations. The advantage of this methodology is that a library may contain thousands of logic functions and not hundreds as in standard cells.

Thus, the logic synthesis tool is able to optimize the circuit with a very large range of logic functions available in the library. After the logic synthesis, an additional step optimize independently transistors according to the loading applied to each gate.

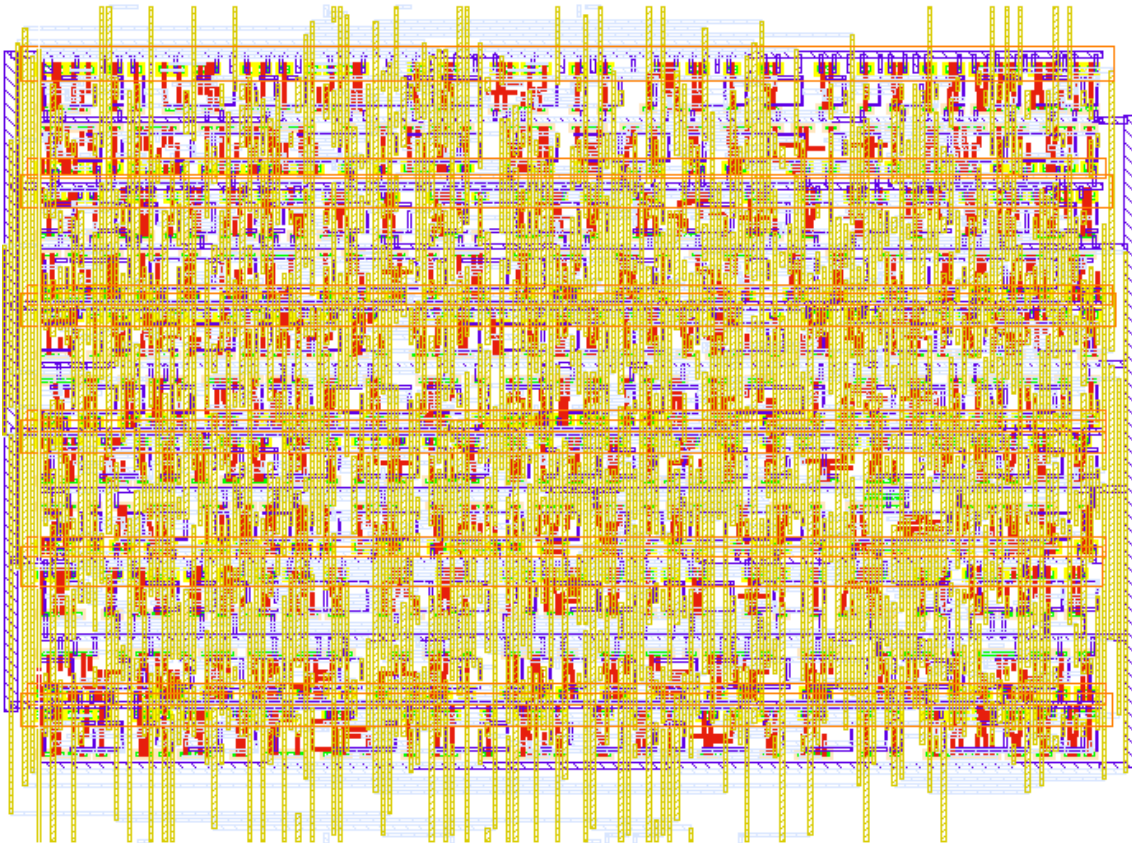
This transistor optimization allows to optimize critical paths while gates in non-critical paths have minimum sizes to reduce the power consumption. Besides, techniques such as the gate length biasing are easily employed in the design process to reduce also the power leakage.

The layout is only generated after the transistor level optimization with computationally efficient algorithms.

The proposed design flow allows designers to mitigate the timing closure problem. Results show that this methodology is very promising. Comparisons between the transistor level methodology and the standard cell approach show interesting results where



(a) C432



(b) C1908

Figure 2.13: Some layouts of circuits generated by the proposed transistor level design flow.

our methodology presented around 11% of delay improvement and more than 30% power savings.

3 ALGORITHMS FOR TRANSISTOR LEVEL AUTOMATIC LAYOUT GENERATION

3.1 Introduction

The automation of layout reduces the design time due to rapid synthesis and enables the designer to deal with a great range of challenges emerging in new process technologies. New technologies challenges require additional functionality as performance-driven placement, area-efficient placement of substrate and well ties, performance-driven detailed routing and layout compaction with preference to critical nets [GMD⁺97].

Old physical synthesis algorithms usually present an exponential complexity and cannot lead with these challenges. Algorithms proposed in the last years are focused in attempting linear complexity or hierarchical methodologies. Hierarchical models can be a solution because they reduce the runtime of exponential algorithms to acceptable levels.

A transistor level design flow is presented in Chapter A.2. Algorithms work with a huge number of variables and careful attention must be taken concerning the complexity and memory usage efficiency.

Algorithms used in the transistor level design flow are presented in this chapter. These algorithms include transistor placement & routing, and compaction. A brief discussion on the algorithm classes and state-of-the-art physical synthesis algorithms are also introduced.

3.2 An Overview of the Algorithm Classes

The synthesis of layout has been widely explored in academic and commercial researches. These algorithms present different characteristics and the resulting efficiency is measured by the runtime and memory usage.

One fundamental point when developing a new algorithm is the tradeoff between the algorithm complexity and the quality of results. According the problem complexity, algorithms with low complexity may lead to an enough quality of results while complex algorithms present unacceptable runtime efficiency.

On the other hand, complex algorithms may present excellent quality of results in a problem with a small number of variables, while runtime efficient algorithms may lead to simplistic results.

The physical synthesis algorithms can be broadly divided into two classes:

- *Deterministic* : A deterministic algorithm is based on a model in which no randomness is involved. Deterministic models use mathematical representations of the underlying regularities that are produced by the entities being modeled and generate theoretically perfect data.
- *Stochastic* : Stochastic models use computational elements that represent the entities and the processes by which they interact and create a procedural algorithm to generate realistic data. Stochastic algorithms present certain randomness. This randomness is usually based on fluctuations observed in historical data or nature's phenomena.
- *Mixed Deterministic-stochastic* : Some algorithms include a mixed solution to solve physical synthesis problems. They usually converge to a solution by a set of deterministic iterations. For each iteration, variables are fed by stochastic models.

3.2.1 Deterministic Algorithms

The class of deterministic algorithms include *constructive* and *Analytical* methods.

Constructive methods model the convergence to a solution by placing one element at a time. Many algorithms have been presented to solve physical design problems. The sequence of elements inserted in the model influences the convergence to a solution.

A typical example of the constructive method may be the *force-directed*. A force-directed based algorithm models a system as a graph where the position of the vertices indicates the solution.

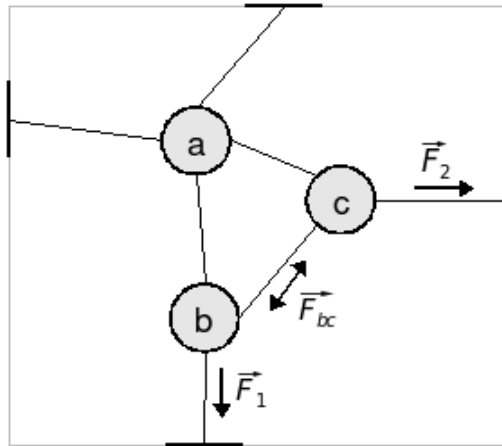


Figure 3.1: A force-directed example.

The force-directed model is represented as a mechanical system of objects connected by springs. The spring behavior is defined by the Hooke's law, which states that the force pushing back the spring is linearly proportional to the distance from its equilibrium length. The Hooke's law is given by

$$F = -kx \quad (3.1)$$

where F is the restoring force exerted by the spring, x is the distance the spring is elongated by, and k is the spring constant or force constant of the spring.

Figure 3.1 exemplifies the force-directed model where only the three forces $\vec{F}_1$, $\vec{F}_2$ and $\vec{F}_{bc}$ are shown. Other forces are omitted. The position of the nodes b and c is a consequence of these forces. Forces $\vec{F}_1$ and $\vec{F}_2$ pull the nodes to the boundaries while force $\vec{F}_{ab}$ controls the attraction between them.

Force-directed models are usually used in algorithms to solve placement problems [WRJS74, MTB00, HCR⁺03, CCS05].

Another constructive algorithm is the placement of elements based on the Euler's path [RTL95, RS99, RS03]. The Euler path algorithm also models the placement as a graph. The idea of the Euler path is to walk on the graph edges where each graph edge is visited once [Wei07b].

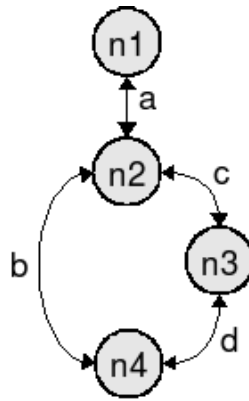


Figure 3.2: An Euler path example.

The Euler path graph is shown in Figure 3.2 where two possible Euler paths could be the sequence of edges a, b, d, c or a, c, d, b .

Analytical methods solve a given problem at once as a system of equations for all components. Analytical formulations are solved by mathematical programming such as *linear programming* (LP) and *quadratic programming* (QP). Linear programming in which variables assume on integer values is known as *integer linear programming* (ILP) or only *integer programming* (IP).

Linear programming is the optimization of an outcome based on some set of constraints using a linear mathematical model [NW07]. Thus, a LP can be used to solve analytical problems if the problem may be formulated as

$$\begin{aligned}
 &\text{minimize} && c^T x \\
 &\text{subject to} && Ax \leq b \\
 &\text{where} && x \geq 0 \\
 &&& L \leq x \leq U
 \end{aligned}$$

where $x \in \mathbb{R}^n$ is a vector of variables that are continuous real numbers. $c^T x$ is the objective function and is represented as $f(x_1, x_1, \dots, x_n) = a_1 x_1 + a_2 x_2 + \dots + a_n x_n + b$,

and $Ax \leq b$ represents the set of constraints. L and U are vectors of lower and upper bounds on the variables.

The most known linear programming solver is the `lp_solve` [Ber06]. `lp_solve` is able to minimize an objective function taking into account linear equalities and inequalities.

Linear programming has been used in placement & routing algorithms [GH00, BC05, BVR06, RC06]. The main characteristic of these works is related to the fast convergence to a solution. Besides, solutions tend to be optimal according to the linear formulations.

Although the excellent characteristics of the linear programming, most part of the problem cannot be represented by a linear objective function.

The quadratic programming is a problem with quadratic objective and linear constraints.

$$\begin{array}{ll} \text{minimize} & q(x) = g^T x + \frac{1}{2} x^T H x \\ \text{subject to} & Ax \geq b \\ \text{where} & L \leq x \leq U \end{array}$$

where $x \in \mathbb{R}^n$ is a vector of variables that are continuous real numbers. $Ax \geq b$ represents the set of constraints. L and U are vectors of lower and upper bounds on the variables.

Some placement algorithms such as the proposed in [AC99, CAL02, KSJA91, VC04] are based on quadratic programming. Thus, the placement problem is described in a mathematical language. Once the formulation converges to a solution, the result is the position of the transistors in the layout.

3.2.2 Stochastic Algorithms

Some design problems have a large range of possible solutions. These problems are computational hard or even impossible to be solved. Some stochastic methods may be used to reduce the search space.

Main examples of stochastic methods are *simulated annealing* and *genetic algorithms*. Simulated annealing is analogous to physical annealing process. It basically involves perturbing independent variables by random values while the temperature controls the standard deviation used by the random number generator. Many placement algorithms based on the simulated annealing have been proposed [MG88, Sec88, Hen02, Tay03, HFPR06].

Genetic algorithms use basic principles of biology and emulates the natural process of evolution to find solutions to a problem.

In the genetic algorithm, each solution is represented by a *chromosome*. A chromosome is usually composed by a binary vector where variables formed by one or more bits are described. A population of chromosomes (possible solutions) is then created and genetic operators as mutation and crossover are applied in order to evolve the solutions to better results.

Some applications of the genetic algorithms on the physical design are presented in [BBR02, LAR05b]. A new transistor placement technique is presented in Section 3.6 that consists on a genetic algorithm integrated with analytical programming.

3.3 Goals of Placement Algorithms

Hentschke [Hen02, Hen07] presents an interesting discussion concerning placement algorithms. In this work, it is presented some characteristics about classic placement algorithms and the importance of these algorithms in modern physical designs. Furthermore, a comparison study on the most important placement algorithms is reported in order to present the consequences of these methods in the wire length and occupied area. Some of these placement algorithms are briefly reported in the following, based on the Hentschke's considerations.

Routability is defined by the ability of routing algorithms to route all wires under electrical and topological restrictions. The responsibility by a non-routable circuit is not only due to the router but also the placer. Otherwise, a router will never be able to route a circuit whether cells are not adequately placed. Thus, the placer must be able to estimate the wire length of a circuit. The main estimation techniques are:

- *Semiperimeter*: Find the smaller *bounding box*, including all cells connected by a net and calculates the Semiperimeter (width + height).
- *Complete Graph*: Calculates the distance between each two points of a net.
- *Source to Target*: Every net has at least one source and one target. This method estimates the wire length from an output pin (source) to the input pins (targets) from a net. This estimation method is important when *timing* is considered.
- *Spanning Tree*: Spanning Tree is a tree where one point is connected only with the nearest point of the net.
- *Steiner Tree*: This wire length method is very realistic and close to the results given by the *Maze routing* algorithm. In Steiner Trees, connections can be done not only by two points but between a point and connection.

The wire length of a circuit is important but the balanced distribution of the connections is mandatory to avoid the congestion. The congestion is indirectly minimized with the wire length, but it is not a guarantee that there is no congestion points in the circuit.

Power dissipation is another important factor in VLSI designs. Higher Clock frequencies associated with the device scaling in deep submicron technologies are increasing the power dissipation in modern circuits. Thus, total power reduction techniques must be taken into account when developing new algorithms for physical design.

However, placement oriented on total power dissipation reduction may provoke areas with high power concentration. Furthermore, placement algorithms targeting power consumption must be able not only to reduce the power dissipation, but they must deal with the distribution of power in the whole circuit.

Deep submicron technology process present more challenges to physical designers concerning *timing* of connections. The wire length is not the only factor considered by a placer, also electric problems such as timing and signal integrity must be taken into account when placement is done.

In nowadays technologies, several metal layers are available and connections are each time smaller. The increase of the resistance due to the reduction of the connections widths and the larger capacitance between lines may cause signal degradation and bigger delay. In the worst case, crosstalk may change the logic signal of a net.

The *area* occupied by a circuit must be considered by the placement algorithm. In order to deal with this constraint, placement algorithms must be able to generate balanced rows and to manage very well empty spaces. A balanced distribution of cells in the rows is the key to reduce the area occupied by the circuit.

3.4 Goals of Routing Algorithms

One of the most important issues imposed by recent technologies is related to circuit wires [Hen07]. Designs are getting bigger while component sizes are becoming dramatically smaller. This *scaling* scenario imposes larger, denser and more complex wiring nets.

Considering *timing* issues, the amount of delay of the logic is being reduced in comparison with the interconnection delay. Interconnection delay is responsible for more than 50% of a circuit delay in submicron technologies.

Considering the *power consumption*, which is strongly affected by the capacitance of circuit nodes. Large wires represents considerably large capacitance to be charged or discharged.

In other words, routability, timing, power and manufacturability are strongly affected by interconnect complexity of a design. In order to cope with wire related problems, the effort of reducing wire length is a very relevant issue on physical design research. Shorter wires are faster, dissipate less power, lead to less complex wiring networks affecting routability and manufacturability.

3.5 State-of-the-art Algorithms for transistor placement and routing

Many algorithms have been presented in the literature during the last decades. These algorithms are related with *Placement & Routing* and *Compaction* techniques and aims to face new technology process challenges. Some of them are discussed in this Section.

3.5.1 Transistor Placement Using an O-tree Algorithm

A non-slicing floorplaning based on ordered trees (O-tree) is presented in [GCY99]. The O-tree placement algorithm presented in this paper is characterized by representing the placement of blocks using an ordered tree structure.

A n -node O-tree is a tree with $n+1$ nodes and encoded by a 2-bit string $\mathbf{T}$ to identify the branching structure of the tree, and a permutation π as the labels of the n nodes. The bit string $\mathbf{T}$ is a realization of the tree structure. We write a '0' for a traversal which descends an edge and a '1' when it subsequently ascends that edge in tree. The permutation π is the label sequence when we traverse the tree in depth-first search order. The first element in permutation π is the root of the tree.

Given the tree shown in Figure 3.3, it can be represented by (00110100011011, ad-bcegf). Thus, starting from the root, we visit node a first and record a bit '0' to $\mathbf{T}$ and

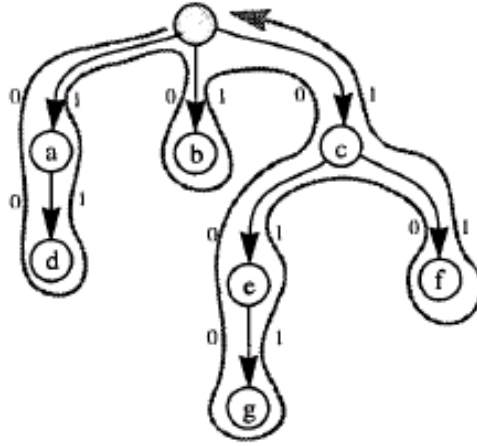


Figure 3.3: Example of a 8-node tree.

a label 'a' to π . Then we visit node d and record a bit '0' to $\mathbf{T}$ and a label 'd' to π . On the way back to the root from nodes d and a , we record two bits "11" to $\mathbf{T}$. Then we visit subtrees b and c in sequence, and record the remaining of $\mathbf{T}$ and π respectively.

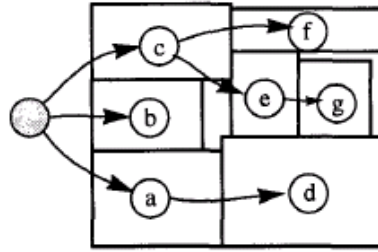


Figure 3.4: The O-tree representation placement.

The root of the O-tree represents the left boundary. Thus, the x-coordinate is given by

$$x_j = x_i + w_i \quad (3.2)$$

where the element i is the parent of the element j , w_i is the width of the element i , x_i and x_j are the left most position of the elements i and j , respectively. In the O-tree placement, x_{root} and w_{root} must be considered as ZERO. Figure 3.4 shows the placement which is represented by the horizontal O-tree in Figure 3.3.

The permutation x determines the vertical position of the component when two blocks have proper overlap in their x-coordinate projections. For each element B_i , let $\Psi(i)$ be the set of B_k with its order lower than B_i in permutation π and interval $(x_k, x_k + W_k)$ overlaps interval $(x_i, x_i + w_i)$ by a non-zero length. If $\Psi(i)$ is non-empty, we have

$$x_i = \max_{k \in \Psi(i)} y_k + k_k \quad (3.3)$$

otherwise

$$y_i = 0 \quad (3.4)$$

The definition of **LB-compact** placement is used to guarantee the best placement according to the O-tree representation. Thus, for a given tree encoding, the algorithm returns the most possible compacted placement. When all elements are placed in the left side, the solution is considered as **L-compact**. A solution is **B-compact** whether all elements are placed in the bottom.

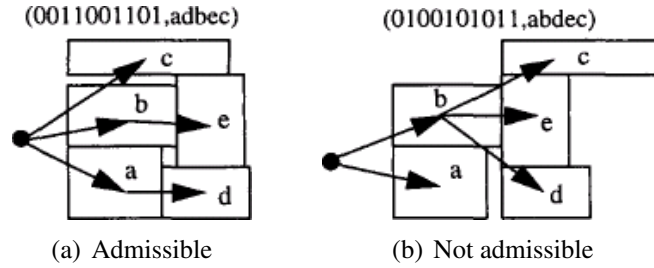


Figure 3.5: Admissible o-tree.

A solution is considered as admissible whenever the resulting placement is a **LB-compact** placement, being both L-compact and B-compact. Figure 3.5(a) illustrates an admissible solution whose elements are placed on the left and the bottom boundaries. Figure 3.5(b) shows an example of not admissible placement solution due to the x-coordinate of the transistors *c* and *d*.

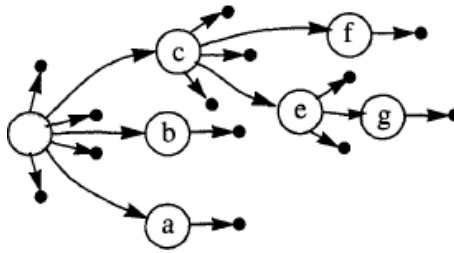


Figure 3.6: Possible insertion position of a external node.

Given an initial O-tree, as shown in Figure 3.3, a new placement configuration can be generated by deleting a component from the O-tree and placing it in another insertion position. For n elements, there is $2n-1$ possible perturbed positions. If the element *d* from the O-tree is chosen to be deleted and permuted, then the possible insertion positions for this element are shown in Figure 3.6. In order to simplify the algorithm, the insertion positions are considered only at the external nodes of the tree.

An automatic datapath placer is presented in [SS01] whose transistors placement are encoded following the O-tree representation presented in [GCY99]. Some modifications were done in order to deal with datapath tile placement characteristics. The main characteristics are the ability to handle fixed tile width, placement on the reflection line (essential to connect adjacent tiles) and non-rectangular shapes.

Figure 3.7 shows the basic idea of the datapath structure. In Figure 3.7(a), the regular structure of a datapath is shown where tiles are placed side-by-side in order to generate the

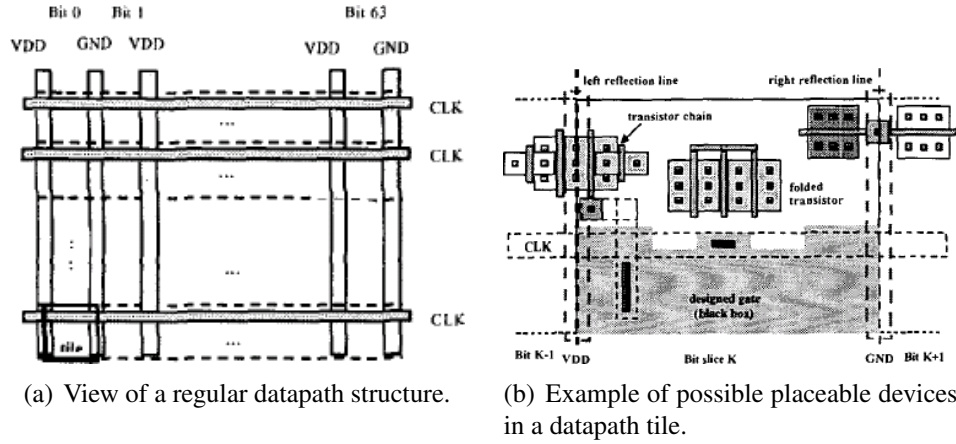


Figure 3.7: Datapath tile placement.

layout. Clock and Supply lines are distributed by the circuit in such a way they connect every tile. Figure 3.7(b) illustrates important characteristics of the datapath generation, where supply lines, the clock line and transistor connections are shared between adjacent tiles.

3.5.2 Transistor Placement with Symmetry Constraints

In high-performance analog circuits, it is often required that groups of devices are placed symmetrically with respect to one or several axis to match the layout-induced parasitics in the two halves. Failure to match these parasitics in differential analog circuits can lead to higher offset voltages and degraded power-supply rejection ratio [BMK04].

Binary trees were used in [Bal00, Bal01, BMK02, BMK04], instead of O-trees, for representing the transistor placement in analog design at device level. Results presented in these works shown that binary trees can efficiently represent the placement of transistors with smaller complexity and saving CPU time.

An important characteristic of these works is related to the possibility to deal with symmetric transistor placement. Symmetry constraints were inserted into the tree encoding, where a subset of tree representations called *symmetric-feasible* is taken into account during the search of the solution space. In these works, simulated annealing is used to explore the solution space.

Figure 3.8 shows the placement of elements with symmetric constraints. In Figure 3.8(a), an example of symmetric placement is illustrated where the pair of transistors (B,H) and (E,F) are symmetrically placed. The O-tree and binary tree encodings are presented in Figure 3.8(b) and 3.8(c), respectively.

Considering the placement of two transistors i and j , symmetry constraints are given by

$$|x_{symAxyS} - x_i + w_i| = |x_{symAxyS} - x_j| \quad (3.5)$$

and

$$y_i = y_j \quad (3.6)$$

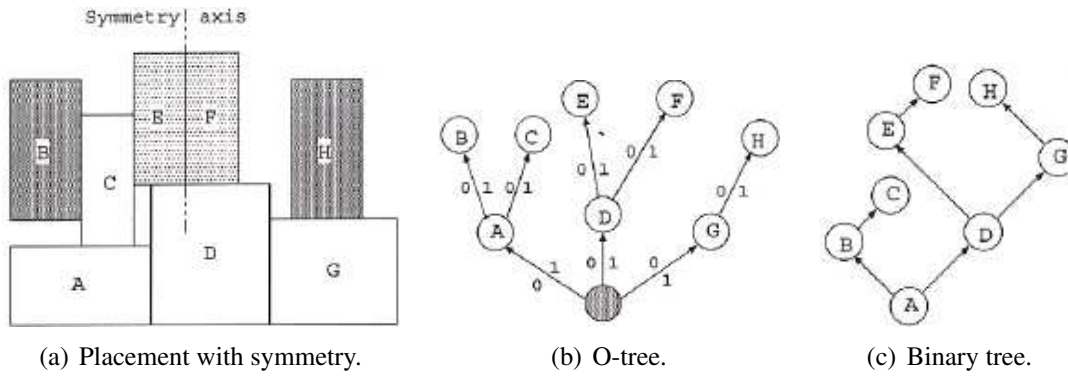


Figure 3.8: Placement with symmetry group and different encodings.

where the element i is always in the left side of the element j and $x_{symAxis}$ is the x-coordinate of the axis of symmetry.

3.5.3 Non-complementary Transistor Placement

There is an increasing need in modern VLSI designs for circuits implemented in high-performance logic families such as Cascode Voltage Switch Logic (CVSL), Pass Transistor Logic (PTL), and domino CMOS. Circuits designed in these non-complementary logic families can be highly irregular, with complex diffusion sharing and nontrivial routing. Traditional digital cell layout synthesis tools derived from the highly stylized “functional cell” style break down when confronted to such circuit topologies. These cells require a full-custom, two-dimensional layout style, which currently requires skilled manual design.

A methodology for the synthesis of such non-complementary digital cell layouts is presented in [RS03]. The methodology permits the concurrent optimization of transistor chain placement and the ordering of the transistors within these diffusion-sharing chains. The mechanism for supporting this concurrent optimization is the placement of transistor subchains, diffusion-break-free components of the full transistor chains. When a chain is reordered, transistors may move from one subchain (and therefore one placement component) to another. This permits the chain ordering to be optimized for both intra-chain and inter-chain routing. The placement algorithm is combined with third-party routing and compaction tools in order to finish the synthesis process.

The methodology presented by Riepe and Sakallah (Figure 3.9) can be summarized by the following points when defining the cell-level transistor placement and routing problem:

1. The input to the system is a sized transistor netlist, the process design rules and technology parameters, and a description of the cell layout style.
2. Transistor source/drain geometry sharing is encouraged, but is not the primary optimization objective. Obtaining a routed cell of minimum area is the primary objective.

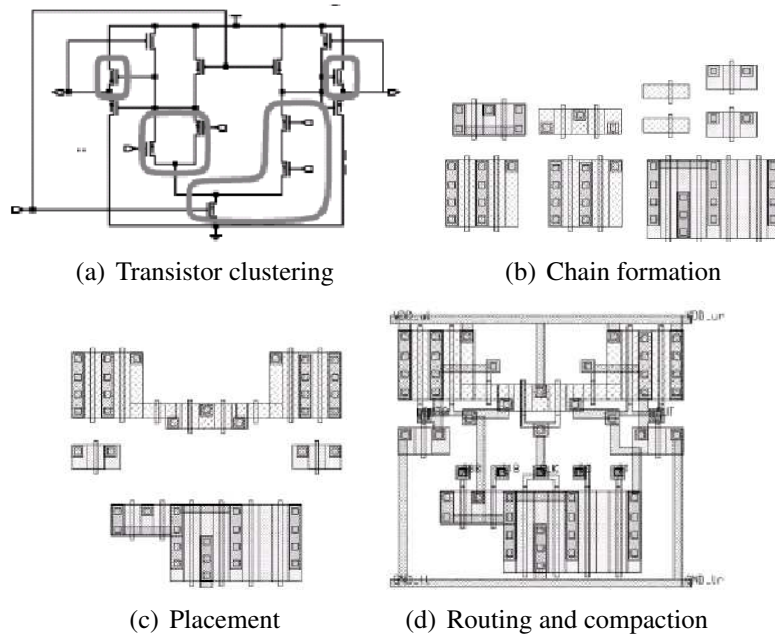


Figure 3.9: Stages of the layout synthesis proposed in [RS03].

3. Individual transistors, or chains of source/drain connected transistors, may be placed in any position or orientation to optimize the objective as long as the design rules and template constraints are satisfied.
4. Routing may be performed in two layers: Polysilicon and first-level metal. There is no preferred direction for routing in either layer.

3.5.4 Transistor Placement & Routing By Using Linear Integer Programming

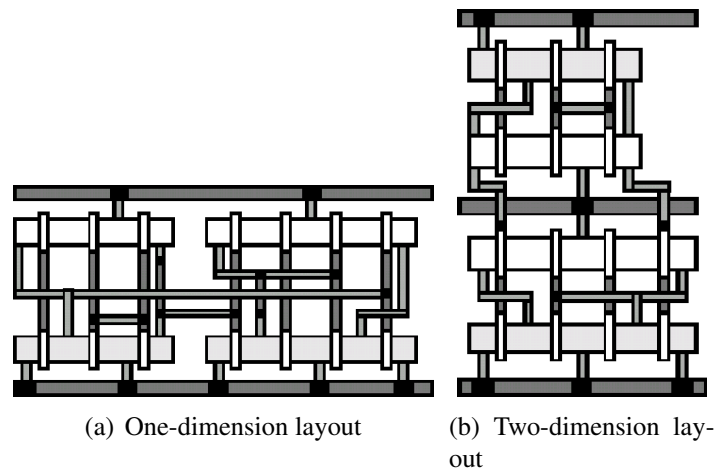


Figure 3.10: The CLIP layout style.

In [GH00], it is presented a technique for the automatic generation of layouts of

CMOS cells in the two-dimensional (2D) style. The technique, CLIP (Cell Layout via Integer Programming) is based on integer-linear programming and solves both width and height minimization problems for 2D cell.

Width minimization is formulated in form that combines factors influencing the 2D cell width in a common problem space: transistor placement, diffusion sharing and vertical inter-row connections. This space is searched in a systematic manner by the branch-and-bound algorithms used in ILP solvers. For height minimization, cell height is modeled based on the horizontal wire density.

The CLIP run time for width minimization is in seconds for circuits with 30 or more transistors. For both height and width optimization, the CLIP is practical for circuits with up to 20 transistors. To extend the algorithms to larger circuits, hierarchical methods are necessary.

Figure 3.10 shows a one-dimension layout in Figure 3.10(a) and, the same two-dimension layout in figure 3.10(b). It is important to note that the three routing horizontal tracks in the two-dimension layout are distributed in the two-dimensional layout.

3.5.5 A Maze Routing Steiner Tree with Effective Critical Sink Optimization

An algorithm for optimized steiner tree generation is presented in [HNJR07]. The algorithm called AMAZE consists on the application of a biasing technique as the key to achieve wire length reduction by maximizing wire sharing. On the other hand, repulsive biasing, path length and sharing factors were introduced to isolate critical paths so that delay to the identified critical sinks is minimized.

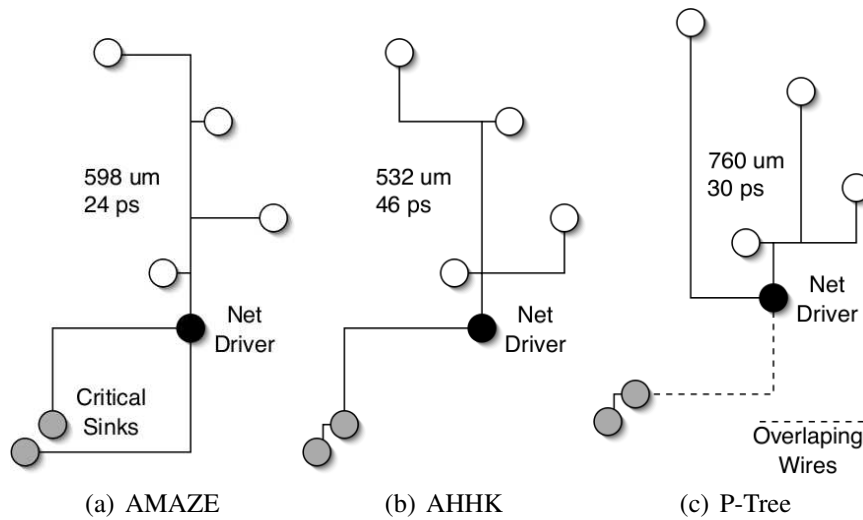


Figure 3.11: Comparison of the best trees generated by (a) AMAZE [HNJR07], (b) AHHK [AHH⁺95] and (c) P-Trees [LCLH96] algorithms.

Figure 3.11 compares nets produced by 3 algorithms. We observe that, with the AMAZE algorithm, the path for the critical sinks has minimum length and minimum sharing, while the rest of the tree is optimized for wire length. The best effort of AHHK [AHH⁺95] algorithm found the minimum path length to all sinks (arborescence), how-

ever it didn't help much for the delay of the critical sinks, since the path is fully shared by both sinks.

With the P-Trees algorithm [LCLH96], among 35 different topologies evaluated, the one that was best for minimizing delay to critical sinks has separated wires to the critical sinks. The P-Trees's drawback that impacted the delay is the fact that the overall wire length is too large (affecting the product of driver resistance per total capacitance).

Among the various topologies generated, the one with better isolation of the critical path failed to provide good wirelength for the rest of the tree. Another drawback is the fact that overlapping wires could be used for Steiner trees and global routing but not for detailed routing.

The AMAZE algorithm outperformed algorithms used in the industry and in the state-of-the-art academic research, such as AHHK by 25%–40% and P-Trees by 1%–30%.

3.5.6 Routing with a Negotiation-based Algorithm

A negotiation-based algorithm for cells routing is presented by Ziesemer in [Jun07]. The negotiation methodology was previously proposed to routing FPGA circuits [ME95]. Ziesemer has used the methodology to generate cells very efficiently.

The algorithm is based on the competition of resources. Thus, nets compete to get congested nodes and nets with more difficulties to find an alternative path have more probability to get the node. Authors guarantee that the negotiation-based algorithm is better than the traditional *rip-up and re-route* algorithms.

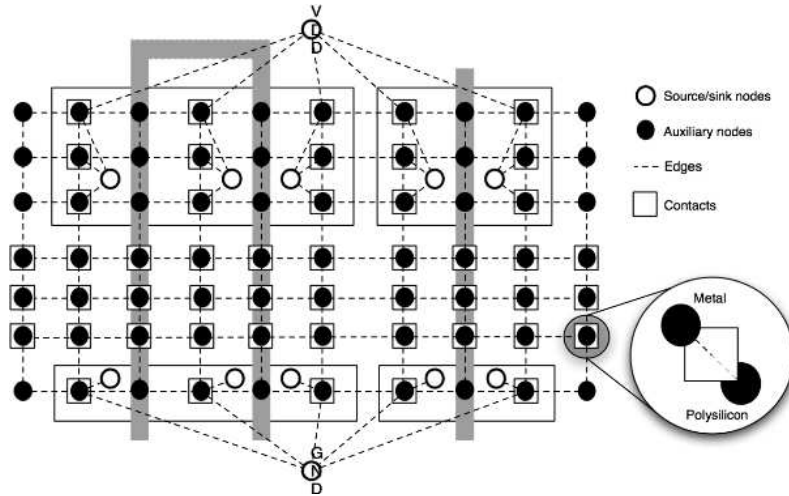


Figure 3.12: Data structure for the negotiation-based algorithm.

Figure 3.12 illustrates the data structure for the negotiation-based algorithm. The layout is represented by a graph where each node has an associated cost.

The cost to use a given node in an iteration is given by

$$C_N = (B_n + H_n) \times P_n \quad (3.7)$$

where B_n is the basis cost for the edge to achieve the node n , H_n is the congestion history of the previously iterations and P_n is the number of paths crossing node n .

To achieve better electrical characteristics of the resulting cell layout, different weights are applied to the graph edges according to its layer and position. Connections in polysilicon mean a bigger cost than in metal and contacts have an even bigger cost. Connections in metal over the transistor gates have an increased cost since they frequently insert additional space in the cell width and therefore must be avoided.

Input/output port connections require additional space to be placed and also there are fixed rules that must be followed. A chain of two or more serial transistors in the same diffusion row can be placed without diffusion contacts between the gates. This allows an area reduction if no ports are placed in the closest track to the transistors. For this reason, better area results are achieved when increasing the routing cost of the graph edges that lead to these ports.

3.6 Transistor Placement Technique Using Genetic Algorithm And Analytical Programming

We propose a new transistor placement technique using genetic algorithm associated to analytical programming in [LAR05b]. The approach presented in this work is basically divided in three phases. First, a classical genetic algorithm is used to generate some parameters concerning transistor orientation and the relationship between them. These parameters are used as placement constraints described in an algebraic modeling language. The second phase consists on solving the placement constraints by a nonlinear solver in order to find the optimal solution according to given constraints. After that, the best solutions are propagated and genetic operators are applied to the solutions.

Algorithm 4 The genetic algorithm.

Require: Set of Transistors T , Population Size N , Number of Iterations I

Ensure: Transistor placement p

```

1:  $P \leftarrow \text{generatePopulation}(T, N)$ ;
2: while  $i < I$  do
3:   for all  $k \in P$  do
4:      $\text{solveConstraints}(k)$ ;
5:      $\text{calculateFitness}(k)$ ;
6:   end for
7:    $P \leftarrow \text{doEvolution}(P)$ ;
8:    $i++$ ;
9: end while
10:  $p \leftarrow \text{getBestSolution}(P)$ ;

```

The pseudo code of the proposed approach is presented in Algorithm 4. An initial set of solutions is generated in the function `generatePopulation( N )` where each chromosome in the population P has a set of constraints about the transistor placement problem. The generation of this initial population is explained in Section 3.6.1.

In function `solveConstraints( k )`, the parameters of the chromosome k are converted to an algebraic modeling language and the placement problem is solved.

The fitness of a chromosome is generated in the function `calculateFitness( k )`

). The fitness of a chromosome is calculated based on the objective function as described in Section 3.6.3.4.

The function `doEvolution( P )` is basically the reproduction of the chromosomes in the population P to generate a new population with better results. In the generation of this new population, operations of elitism, mutation and crossover are applied to the chromosomes in order to propagate the best solutions and to evolve the other chromosomes.

3.6.1 Initial Population Generation

The range of possible solutions in the process of layout generation is related to the number of elements in a cell or macro-block. Moreover, the relation between these elements makes a solution better than others. Thus, some techniques can be used for reducing the number of elements and, consequently, decreasing the complexity of the layout generation problem.

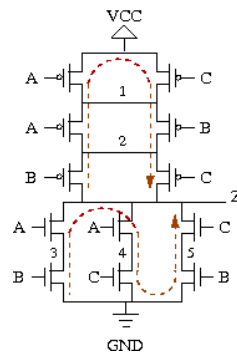


Figure 3.13: Euler path example.

Transistor chaining is a technique that consists of grouping transistors when their drain/source diffusions can be shared. Figure 3.13 illustrates the transistor chaining generation where the *Euler* path is searched to PMOS and NMOS transistors. Dashed lines illustrate *Euler* paths in which a chain of transistors is performed based on the sharing of the source/drain diffusion areas. In this example, the two transistors chains are $(Z, B, 2, A, 1, A, VCC, C, 1, B, 2, C, Z)$ to PMOS transistors and $(GND, B, 3, A, Z, A, 4, C, GND, B, 5, C, Z)$ to NMOS transistors.

It is clear that many solutions can be found to these set of transistors. In the approach proposed in this work, an *Eulerian* graph is used in order to generate the N solutions related to the initial population. Transistor chainings are randomly chosen to be used in the genetic algorithm.

3.6.2 The Placement Parameters

Each chromosome in the genetic algorithm is a set of parameters used in the placement constraints. Parameters used in transistor placement are basically the description of transistors orientation and the relationship between these transistors. Transistor orientation means whether a transistor must be placed horizontally or vertically and where the

drain/source contacts are located, while the relationship between transistors is the relative placement of a transistor in relation to each other transistor.

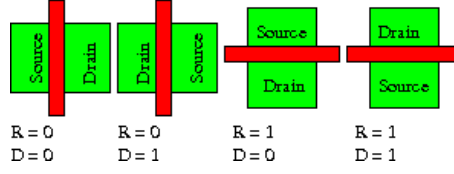


Figure 3.14: Transistor orientation constraints.

Figure 3.14 illustrates the orientation constraints R and D . The parameter R represents the orientation of the transistors. $R = 0$ indicates that the transistor must be placed horizontally and $R = 1$ means that the transistor must be placed vertically.

The parameter D indicates where drain/source diffusion areas are located. $D = 0$ means that the transistor source area is located in the left/top and $D = 1$ means that the drain area is located in the left/top of the transistor.

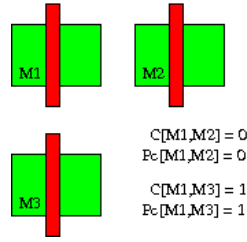


Figure 3.15: Transistor behavior constraints.

The relationship between transistors is shown in Figure 3.15. The parameters C and P_C are used to describe these relationship. C indicates whether the placement constraints are related to horizontal or vertical coordinates and P_C represents the relative position of these transistors.

Taking as example the transistors $M1$, $M2$ and $M3$ illustrated in Figure 3.15, $C[M1, M2] = 0$ means that the transistors $M1$ and $M2$ are placed side by side horizontally and $P_C[M1, M2] = 0$ indicates that the transistor $M1$ is placed in the left side of $M2$. In other words, $X_{M1} < X_{M2}$ and there is no requirements to coordinate Y .

The same idea is used to $C[M1, M3] = 1$ and $P_C[M1, M3] = 1$. In this case $Y_{M1} > Y_{M3}$ and any horizontal constraint is applied. Table 3.1 shows the possible constraints resulting of the parameters C and P_C .

Based on these parameters, each chromosome is a binary vector containing information about orientation and relationship between transistors. The size of a chromosome is given by Equation 3.8:

$$L_{chrom} = T * 2 + \sum_{i=1}^{T-1} i * 2 \quad (3.8)$$

where T is the number of transistors. The first part of the equation 3.8 is related to parameters R and D , and the second part is related to parameters C and P_C .

Table 3.1: Parameters to the transistors relationship.

Parameters		Constraints	
C	Pc	Horizontal	Vertical
0	0	$X_1 < X_2$	—
0	1	$X_1 > X_2$	—
1	0	—	$Y_1 < Y_2$
1	1	—	$Y_1 > Y_2$

3.6.3 The Mathematical Modeling

Once the parameters are defined in the chromosomes, they can be applied in an algebraic modeling in order to obtain the optimal placement solution to given parameters and constraints. The main idea of this approach is to use a nonlinear solver to find the solution to the placement of transistors.

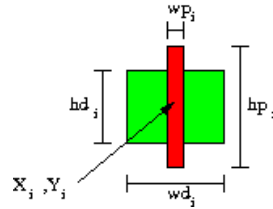


Figure 3.16: Width and height transistors parameters.

Figure 3.16 shows width and height parameters used in the placement constraints. For each transistor $i \in T$, the parameters wd_i and hd_i are the width and height of the diffusion area while wp_i and hp_i are the parameters for the polysilicon area. Besides, three integer parameters $drain_i$, $source_i$ and $gate_i$ represent the connections of the transistors and the parameter $type_i$ is also used to indicate PMOS and NMOS transistors.

The variables X_i and Y_i are the central coordinates of the transistor i . Their values are given by the minimization of the objective function. The goal of the used objective function is to find the optimal X_i and Y_i by the minimization of the wire lengths. The specification of the objective function is given in more details in Section 3.6.3.4.

The constraints are divided in three groups: 1) Boundary Constraints, 2) Neighborhood Constraints and 3) Connections Constraints.

3.6.3.1 Boundary Constraints

The layout of standard-cells and macro-blocks is usually structured in rows. In these structures, layout boundaries must be regular in order to allow the connection between adjacent cells at the moment of circuit generation. Figure 3.17 illustrates the boundaries in a row-based layout.

Regions for PMOS and NMOS transistors can be determined by the implant areas and boundary constraints can be formulated according to the edges of these areas. Thus, the

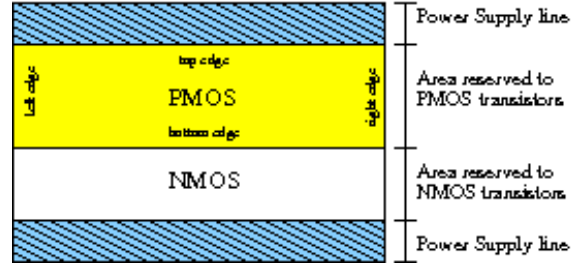


Figure 3.17: A row-based boundary representation.

boundary constraints are given by

$$B_{left} + \Delta_x + \frac{1}{2}W_i \leq X_i \leq B_{right} - \Delta_x - \frac{1}{2}W_i \quad (3.9)$$

$$B_{bottom} + \Delta_y + \frac{1}{2}H_i \leq Y_i \leq B_{top} - \Delta_y - \frac{1}{2}H_i \quad (3.10)$$

where B_{left} , B_{right} , B_{bottom} and B_{top} are the edges of the placement region, Δ_x and Δ_y are the minimal distances from the transistor i to the boundaries, and W_i and H_i are the width and height of the transistor.

3.6.3.2 Neighborhood Constraints

Neighborhood constraints are related to the possibility to connect transistors. These constraints are separated in categories and they are responsible to give the correct distance between two adjacent transistors.

In order to verify the possibility of connection between transistors, the variables *left*, *right*, *top* and *bottom* are used. They are given by

$$left_i = (D_i + R_i * D_i) * drain_i + (1 - D_i + R_i * D_i) * source_i + R_i * gate_i \quad (3.11)$$

$$right_i = (1 - D_i + R_i * D_i) * drain_i + (D_i + R_i * D_i) * source_i + R_i * gate_i \quad (3.12)$$

$$top_i = (R_i - R_i * D_i) * source_i + (R_i * (1 - D_i)) * drain_i + (1 + R_i) * gate_i \quad (3.13)$$

$$bottom_i = (R_i - R_i * D_i) * drain_i + (R_i * (1 - D_i)) * source_i + (1 + R_i) * gate_i \quad (3.14)$$

where $i \in T$, R_i and D_i are the parameters given by the current chromosome. $drain_i$, $source_i$ and $gate_i$ are integer parameters related to the list of connections C .

Considering κ_c the number of points of the connection c and assuming that $c \in C$, it is possible to know when two transistors are connected in series or parallel. Thus, two transistors are in series whenever $\kappa_c = 2$. In all other cases the transistors are in parallel or they are not connected.

From the definition of these variables, it is possible to understand how the neighborhood constraints are formulated. Figure 3.18 illustrates every neighborhood possibility to the horizontal placement and Table 3.2 presents the neighborhood constraints where sp is the spacing between polysilicon lines, sdc is the distance between a polysilicon line and a contact, wc is the width of a contact, sd is the spacing of two diffusion areas and sdp is the distance between a polysilicon line and a diffusion area.

Table 3.2: Horizontal neighborhood constraints.

#	Figure	Orientation Parameters		κ_c	Situation	Constraint
		R_i	R_j			
1	3.18(a)	0	0	$= 2$	$right_i = left_j$	$X_j - X_i \geq \frac{1}{2}wp_i + sp + \frac{1}{2}wp_j$
2	3.18(b)	0	0	$\neq 2$	$right_i = left_j$	$X_j - X_i \geq \frac{1}{2}wp_i + 2 \times spc + wc + \frac{1}{2}wp_j$
3	3.18(c)	0	0	$\times$	$right_i \neq left_j$	$X_j - X_i \geq \frac{1}{2}wp_i + sd + \frac{1}{2}wp_j$
4	3.18(d)	1	0	$\times$	$\times$	$X_j - X_i \geq \frac{1}{2}hp_i + sd + \frac{1}{2}wd_j$
5	3.18(e)	0	1	$\times$	$\times$	$X_j - X_i \geq \frac{1}{2}wd_i + sd + \frac{1}{2}hp_j$
6	3.18(f)	1	1	$\times$	$right_i = left_j$	$X_j - X_i \geq \frac{1}{2}hd_i + \frac{1}{2}hd_j$
7	3.18(g)	1	1	$\times$		$X_j - X_i \geq \frac{1}{2}hp_i + sp + \frac{1}{2}hp_j$

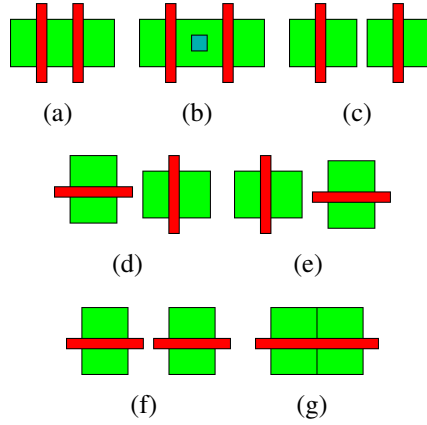


Figure 3.18: Horizontal neighborhood.

Neighborhood constraints are separated in categories with the effort to deal with every possible relationship between two transistors. Only horizontal constraints are discussed here but similar equations are used vertically.

Seven different constraints are shown in Table 3.2. In the case of $R_i = 0$ and $R_j = 0$, equation 1 treats situations where transistors are in series, equation 2 deals with parallel transistors and equation 3 takes situations where transistors are not connected.

The equation 3 and 4 treat situations where there are different transistor orientation parameters ($R_i \neq R_j$). In these cases, the sharing of diffusion areas is impossible.

When transistors are placed vertically ($R_i = 1$ and $R_j = 1$), the connection between two transistors is possible only if $top_i = top_j$, $right_i = left_j$ and $bottom_i = bottom_j$ (Equation 6). Equation 7 takes all other cases to $R_i = 1$ and $R_j = 1$, in which the connection between transistors cannot be done.

3.6.3.3 Connection Constraints

Let n be the number of connections and m the number of transistors, the position to gate, drain and source can be inserted in matrix notation to the horizontal and vertical coordinates, Q_x and Q_y . Thus, Q_x and Q_y are $n \times m$ matrices where the coordinates X

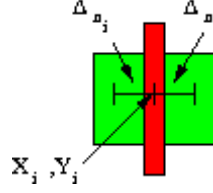


Figure 3.19: The Δ_{n_i} representation.

and Y of the nets are given by

$$Q_x(\text{drain}_i, i) = (1-R_i)*D_i*(X_i-\Delta_{n_i})+(1-R_i)*(1-D_i)*(X_i+\Delta_{n_i})+R_i*X_i \quad (3.15)$$

$$Q_x(\text{source}_i, i) = (1-R_i)*(1-D_i)*(X_i-\Delta_{n_i})+(1-R_i)*D_i*(X_i+\Delta_{n_i})+R_i*X_i \quad (3.16)$$

$$Q_x(\text{gate}_i, i) = X_i \quad (3.17)$$

where $i \in T$ and Δ_{n_i} is the distance from the center of the transistor to the point where the connection is located as shown in Figure 3.19. The matrix to vertical coordinates Q_y is composed based on the same idea.

3.6.3.4 The Objective Function

The goal of the proposed technique is to reduce the wire length connecting the transistors. Thus, the objective function is based on the connection constraints and it is obtained by

$$OBJ : \min \left\{ \sum_{c \in C} W_c * S(c) \right\} \quad (3.18)$$

where $S(c)$ is the half perimeter wire length and W_c is the weight of the connection c . The wire length of a connection c is calculated by the coordinates of the points of a net in the matrices Q_x and Q_y . Then, $S(c)$ is given by

$$S(c) = \sum_{i \in T, j \in T, i \neq j} HP(c, i, j) * I(c, i) * I(c, j) \quad (3.19)$$

and

$$HP(c, i, j) = |Q_x(c, i) - Q_x(c, j)| + |Q_y(c, i) - Q_y(c, j)| \quad (3.20)$$

where $I(c, i)$ are binary values indicating whether the wire c is connected to the transistor i . The same principle is used to $I(c, j)$ with the connection c and the transistor j .

3.6.4 Obtained Results

Figure 3.20 shows the placement of two cells using the proposed algorithm. The transistor placement of an OR2 gate is shown in Figure 3.20(a) and the placement of an AOI222 is shown in Figure 3.20(b).

Table 3.3 shows some results of the comparison between the proposed technique and a pure Eulerian placement algorithm used in [LDGR03]. Results show that the proposed technique deals with the transistor placement problem. The area gain is around 4.5 % .

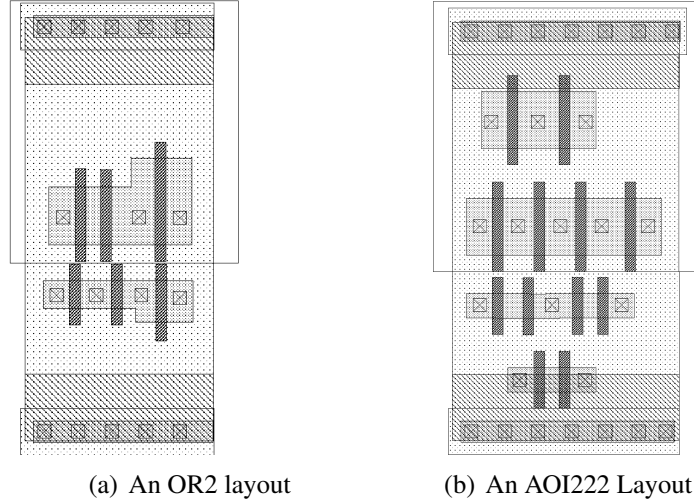


Figure 3.20: Preliminary placement examples.

Table 3.3: Placement results.

Cell Name	Area (μm)		Gain (%)	Execution Time
	[LDGR03]	Proposed		
NOR2	7.9	8.1	-3	2s
OR2	10.8	10.9	-1	1m 15s
AOI22	13.6	13.0	4	4m 10s
AOI222	19.4	17.9	8	18m 30s
Full Adder	52.3	45.1	13	3h 15m

The drawback of this technique is the execution run time. While a pure Eulerian algorithm executes the placement task very quickly, the proposed technique take hours in some cases to solve the placement problem. As the genetic algorithm works with random information, the execution time presented in Table 3.3 is the average time of at least 5 executions of each cell.

3.7 Compaction and Layout Optimization using Linear Programming

Linear programming has been used to solve many design problems such as placement [WLS05], routing [BC05] or even complete cell generation [GH00]. However, describing a problem using only linear constraints may impose undesired restrictions to the design and non optimal resulting layout.

Layout compaction and migration between technologies are applications where the linear programming is very well accepted. In this section, a layout compaction algorithm using linear programming is presented. In order to place each polygon in the layout, its coordinates and technology rules are represented in form of linear equalities and inequalities. The generic linear *LPSolve* [Ber06] is called to solve the compaction constraints.

In order to understand how the linear programming compaction works, it is necessary to have in mind how the constraints are generated. When placement and routing are performed, the relative information of each polygon position is stored in the data structure. Thus, it is possible to know whether a polygon is placed on the left side or on the right side in relation to another polygon.

Before generating linear constraints, the whole layout generation process is free of technology. Transistor placement and routing algorithms do not need any information about the technology rules. The only needed information is which polygon is on the left/right (top/bottom in the case of vertical compaction) side of another polygon. In other words, the relative position of each polygon is known.

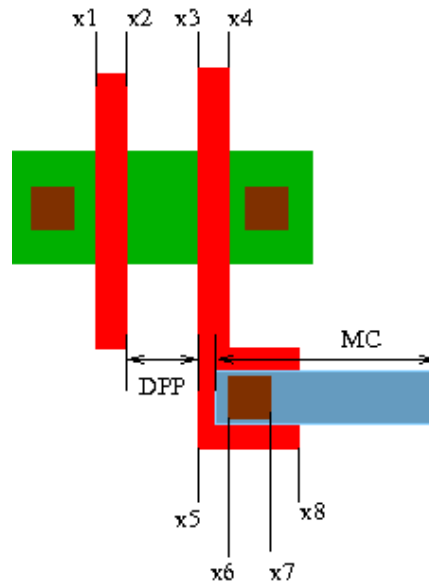


Figure 3.21: Example of layout.

Once the relative position of the polygons is known, it is possible to describe the relation between these polygons as the linear programming. Figure 3.21 illustrates an example of layout, in which two transistors are placed side by side. Constraints related to these polygons can be described as shown in Algorithm 5.

Algorithm 5 An example of layout representation in linear equalities and inequalities.

- 1: $x_2 - x_1 = WP$
 - 2: $x_4 - x_3 = WP$
 - 3: $x_3 - x_2 \geq SP$
 - 4: $x_7 - x_6 = WC$
 - 5: $x_6 - x_5 \geq EPC$
 - 6: $x_8 - x_7 \geq EPC$
 - 7: $x_5 \leq x_3$
 - 8: $x_8 \geq x_4$
-

WP is the polysilicon width and SP is the spacing between two adjacent polysilicon lines. x_1 and x_2 are left and right horizontal coordinates of the left polysilicon line. x_3

and x_4 are the left and right coordinates of the right polysilicon line. The first and second lines of this example shows how the minimum transistor width is guaranteed.

The third line illustrates how the minimum spacing between two polygons can be guaranteed, where the spacing between the polysilicon lines are placed with a spacing equal or greater than the minimum spacing technology rule.

Enclosure constraints are shown in lines 04 to 06. WC is the contact width and EPC is the minimum enclosure of polysilicon over contact. These three constraints impose to the solver a way to find the coordinates for the polygons without violating the technology rules.

The connection between the contact and its enclosure with the polysilicon transistor is depicted by lines 07 and 08. Thus, there is a certain freedom to place the contacts, but the connection is respected.

The main idea of using linear programming in the layout generation is to reduce the area occupied by the circuit. Besides, some important aspects are taken into account in order to generate optimized layout. The main aspects are:

- *Stacked transistors* : Stacked transistor are known to form a resistive path. The distance between two transistors must be the minimum possible in order to reduce the area and perimeter of the transistors drain/source.
- *Polysilicon and diffusion lines* : The resistance of polysilicon and diffusion lines is very high. For this reason, small connections in these layers is mandatory.
- *Transistor active region* : As stacked transistors, active regions must be reduced in order to reduce the area and perimeter of the transistor drain/source regions.
- *Metal lines* : The resistance in metal connections is smaller than polysilicon/diffusion ones, but the reduction of these lines is also important.

The example shown in the Algorithm 5 can be described as follows in order to exemplify how this characteristics are taken into account in the linear programming. In this example, the spacing between the transistors DPP is inserted in the objective function aiming at reducing the distance between them.

Algorithm 6 An example of layout representation in linear equalities and inequalities with an objective function.

- 1: $\min : DPP$
 - 2: $x_2 - x_1 = WP$
 - 3: $x_4 - x_3 = WP$
 - 4: $x_3 - x_2 = DPP$
 - 5: $DPP \geq SP$
 - 6: $x_7 - x_6 = WC$
 - 7: $x_6 - x_5 \geq EPC$
 - 8: $x_8 - x_7 \geq EPC$
 - 9: $x_5 \leq x_3$
 - 10: $x_8 \geq x_4$
-

Each one of the aspects shown in the Algorithm 6 can be considered in linear programming by applying costs to the constraints and inserting them in the objective function. Constraint costs are directly related to the priority of a constraint. For example, the spacing between stacked transistors is very important due to the resistance of the active region of a transistor. A connection in metal has smaller priority. Thus, to improve the quality of the layout, the objective function can be written as

$$\min : 2DPP + MC$$

where DPP is the distance between two polysilicon lines and MC is the length of a metal line. An objective function described in this form results in better layout. Otherwise, if no costs are used, the spacing between the transistors can be harmed as a result of a small metal connection.

Algorithm 7 A compaction algorithm using linear programming.

Require: Set of polygons B

Ensure: The polygons position P

- 1: $D \leftarrow \text{describePolygonsAsLP}(B);$
 - 2: $\text{applyCosts}(D);$
 - 3: $L \leftarrow \text{callLPSolver}(D);$
 - 4: $P \leftarrow \text{placePolygons}(L);$
-

The complete compaction algorithm is presented in Algorithm 7. Function $\text{describePolygonsAsLP}(B)$ converts the relative position of the polygons to linear equalities and inequalities. Costs are applied to the linear programming in Function $\text{applyCosts}(D)$. After, the formulations are solved in Function $\text{callLPSolver}(D)$ and Function $\text{placePolygons}(L)$ places each polygon in the layout.

Figure 3.22 illustrates the layout compaction without costs assignment and with costs attributed to the polysilicon lines, diffusion areas and metal connection. For this example, diffusion areas have higher cost (4), polysilicon lines have a smaller cost (3) and metal lines have cost equal to 1.

3.8 Overview of the Algorithms Developed in the Punch++

As discussed in Chapter A.2, the runtime is one of the most important issues when generating a whole circuit at a time. There is a tradeoff between the quality of circuits and the performance of algorithms that must be carefully analyzed.

Thus, for the development of the automatic layout generator **Punch++**, we choose to sacrifice some layout aspects such as the occupied area to have computationally efficient algorithms.

Basically, the algorithms in an automatic layout generator are separated in three types: *Placement*, *Routing* and *Compaction*.

The **Euler path algorithm** (Section 3.2.1) is used to place the cells. The algorithm is computationally efficient and gives reasonable results. Transistors are separated concerning the logic functions. Thus, the complexity is reduced because the algorithm deal with only tenths of transistors.

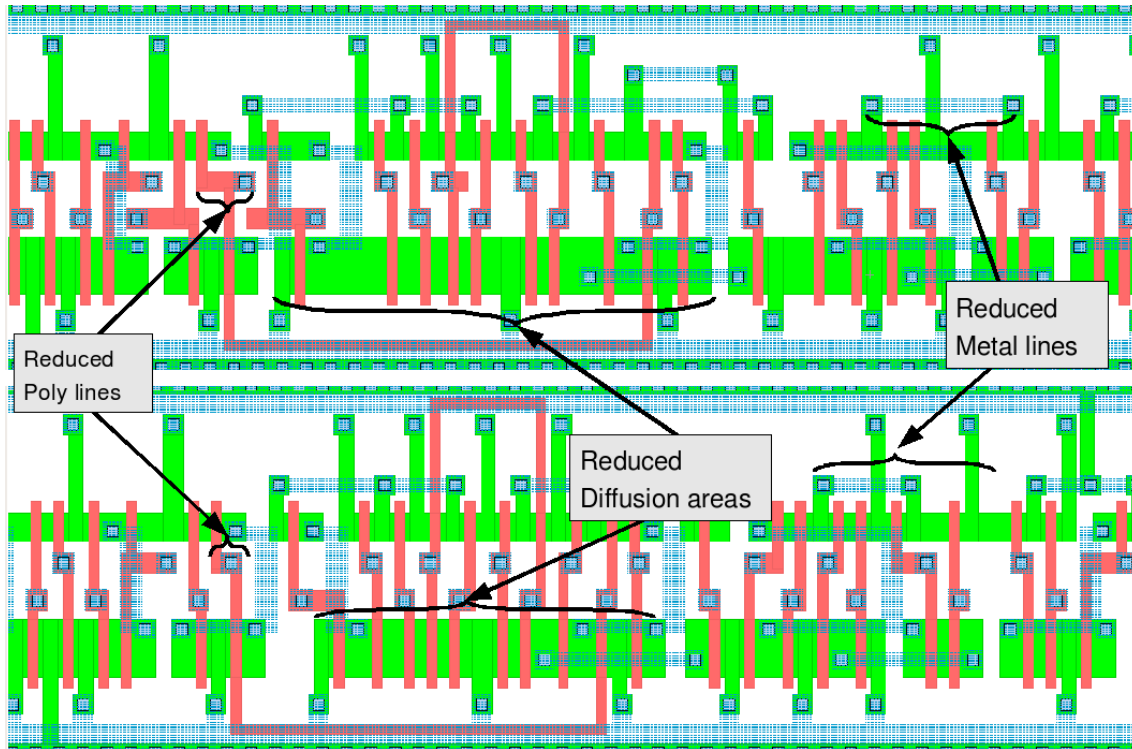


Figure 3.22: An example of layout compaction. The first figure is a layout compacted without cost assignment. Second layout is compacted with costs.

The routing algorithm is based on the **negotiation algorithm** presented in Section 3.5.6. After the placement & routing, the **linear programming-based compaction algorithm** discussed in Section 3.7 is used to compact the layout. The routing and compaction is done in the whole row.

These three algorithms are basically the kernel of the layout generator. They are computationally very efficient and result in a good rapport between layout quality and algorithms performance.

3.9 Conclusion

Physical synthesis algorithms are presented in this chapter. These algorithms are used in placement and routing of transistors aiming at generating whole blocks with up to thousands of transistors.

Algorithms presented in this chapter are the state-of-the-art of the methodologies existing in literature. The main characteristics of the placement, routing and compaction algorithms is to achieve with the recent and growth technology challenges.

A review of the main submicron technology process challenges are shown in Table 3.4.

Table 3.4: Review of the challenges targeted by physical synthesis algorithms.

Runtime	The computational complexity is one of the most important factors when developing new algorithms. Algorithms with linear complexity deal very well with a big number of variables. However, sometimes physical design problems are not easily represented in linear form.
Routability	Routability is defined by the ability of routing algorithms to route all wires under electrical and topological restrictions. Placement and routing algorithms must be able to route the whole circuit and respect design and electrical rules.
Timing	The timing is not only the result of the circuit topology, but also how physical synthesis algorithms deal with the placement, routing and compaction. Algorithm must be able to balance charges in the nets in respect to the driving cell and control the placement to avoid long wires.
Power Consumption	Power consumption is totally related with the transistor placement and routing. There are an enormous set of possibilities of placement and routing possibilities for a given set of transistors. The power consumption may considerably vary among this range of solutions according to the transistors sharing areas and the capacitances associated to the connections.

4 A RADIATION-INDUCED EFFECTS OVERVIEW

4.1 Introduction

Radiation effects such as Total Ionizing Dose (TID), Displacement Damage (DD) and Single Event Effects (SEE) provide aerospace designers' a myriad of challenges for the design of a system [LBM⁺00].

This chapter introduces some effects induced by the radiation. Principles and effects about the total ionizing dose and the displacement damage are briefly discussed. For the purpose of this thesis, single event effects are the focus of this chapter. More specifically, Soft SEE are discussed in details.

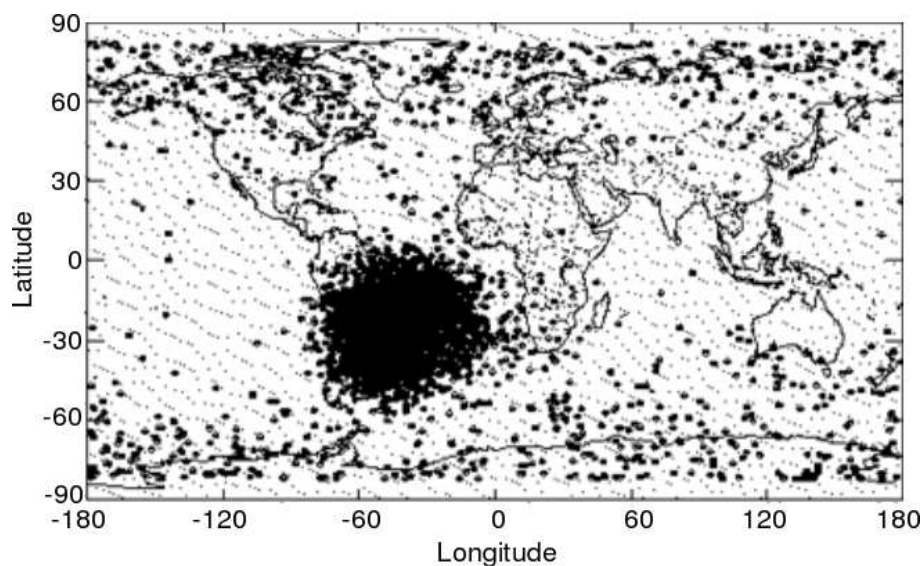


Figure 4.1: Location where Single Event Upsets (SEU) occurred in a spacecraft into a polar orbit of altitude 700km [HSDU⁺90].

Radiation-induced spacecraft anomalies have been known since the Explorer I launch on January 31, 1958, when a Geiger counter put aboard by Van Allen suddenly stopped counting. It turned out that the counter was in fact saturated by an extremely high count rate. This event led to the discovery of the Van Allen belts [MAB⁺03]. The inner belt, beginning at about 1,000 km above the surface of the Earth, contains primarily protons

with energies between $10\text{-}100\text{ MeVcm}^2/\text{mg}$.

The offset between Earth's geographical and magnetic axes causes an asymmetry in the radiation belt above the Atlantic Ocean in the Brazilian coast, allowing the inner belt to reach a minimum altitude of 250 km.

This *South Atlantic Anomaly* (SAA) is important because it occupies a region in which low-orbiting satellites spend as much as 30% of their time. During a solar flare, which can happen anytime, the number of protons suddenly increases by more than a million. Figure 4.1 shows the location where Single Event Upsets (SEU) occurred in a spacecraft into a polar orbit of altitude 700km as presented in [HSDU⁺90].

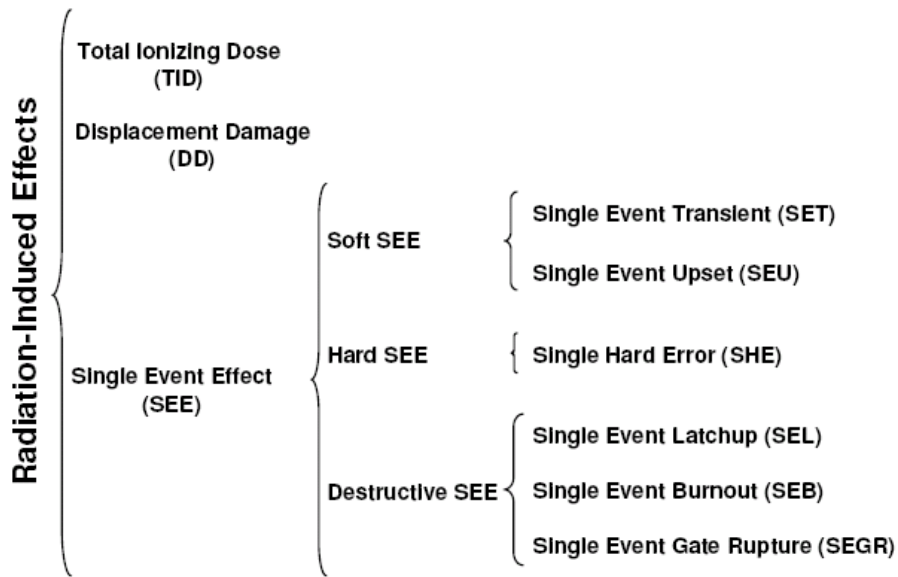


Figure 4.2: Classification of radiation-induced effects [Bas06].

Bastos [Bas06] classifies radiation-induced effects as shown in Figure 4.2. Basically, radiation-induced effects are divided in three categories: Total Ionizing Dose (TID), Displacement Damage (DD) and Single Event Effects (SEE).

Total Ionizing Dose (TID) is due to the degradation as a result of the cumulative energy deposited in the material. This degradation occurs in long term, but the effects are permanent in devices and usually lead to functional failures. Displacement Damage (DD) is also a long term degradation of devices, but is a different physical mechanism. Different from TID and DD, Single Event Effects (SEEs) occur when a single ion strikes the material with sufficient energy in the device to cause a system failure.

4.2 Single Event Effects (SEEs)

Single Event Effects (SEEs) occur when a single ion strikes the material, depositing sufficient energy in the device to cause an error. SEE are divided in Soft and hard errors [LBM⁺00]. Basically, a soft error occurs when a transient pulse or bit-flip causes an error detectable at the device output. Hard errors are physically destructive to the device and cause permanent functional effects.

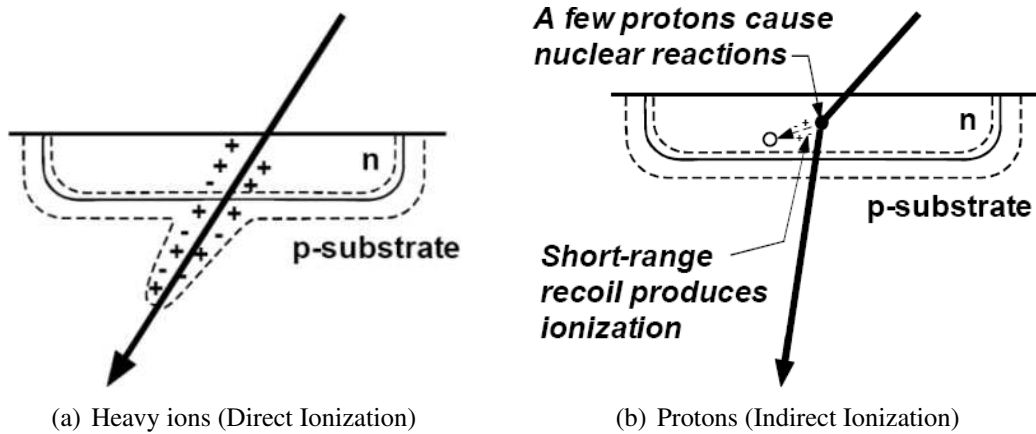


Figure 4.3: Heavy ions and protons striking the silicon device. (a) heavy ion increasing the depletion region (b) *Spallation* caused by a proton or neutron.

4.2.1 Soft SEE

Several types of particles are generated by the sun activity. These particles are classified in two groups: 1) *Charged particles* such as electrons, protons and heavy ions, and 2) *Electromagnetic radiation (Photons)* as x-ray, gamma ray and ultraviolet light [adLK03].

A heavy ion is defined as any ion with atomic number greater than two. A single heavy ion striking a silicon device loses its energy as a result of the production of free electron-hole pairs. The creation of these electron-hole pairs generates a dense ionized track increasing the depletion region as shown in Figure 4.3(a). In other words, heavy ions cause direct ionization within the device.

The direct ionization is the primary charge deposition mechanism for single events. The particle loses its energy during the direct ionization, when crossing the semiconductor material.

The term *linear energy transfer* (LET) is used to describe the energy loss by a particle. The LET of a particle can be easily related with its charge deposition [DM03]. An LET of $97 \text{ MeV} - \text{cm}^2/\text{mg}$ in silicon devices corresponds to a charge deposition of $1 \text{ pC}/\mu\text{m}$.

Lighter particles such as proton, electron and neutrons do not usually produce enough charge by direct ionization to cause a single event. Typically, protons and neutrons cause a transient pulse or a bit-flip through complex nuclear reactions such as emission of alpha and gamma particles or spallation in the vicinity of the sensitive node (Figure 4.3(b)). *Spallation* is a nuclear reaction in which two or more fragments or particles are ejected from the target nucleus.

Any one of these reactions may deposit energy along their paths. The particles created by these reactions are much more heavier than the original neutron or proton. This resulting heavy particles deposit higher charge capable to provoke a single event.

When the energy of a particle is enough to increase the current in a transistor node, a temporary current disturbance is generated. When this disturbance occurs in a combinational logic circuit, the effect is known as single event transient (SET). SETs may lead a system to an unexpected response whether it propagates to a memory element or a primary output (PO) of a circuit.

On the other hand, when the increase of current occurs in storage elements such as latches, flip-flops and RAMs, it may provoke the modification of the logic value stored in these elements. This modification in the value stored in the memory element as a function of a particle is called Single Event Upset (SEU).

4.2.1.1 A Single Event Model

Messenger presents in [Mes82] a fault model aiming at estimating the effects of single events. The model represents the effects of an α -particle striking a device as a double exponential current curve. This curve is obtained by

$$I(t) = \frac{Q}{\tau_\alpha - \tau_\beta} \left(e^{-t/\tau_\alpha} - e^{-t/\tau_\beta} \right) \quad (4.1)$$

where Q is the injected charge and may be positive or negative, τ_α is the collection time constant of the junction and τ_β is the time constant for initially establishing the ion track. τ_α and τ_β are constants and depend on several process-dependent factors.

In the double exponential, the τ_β is responsible to shape the rising of the current pulse and the τ_α shapes the fall time of the curve. The curve presents a fast rise time with smaller τ_β , as well as the fall time is faster with smaller τ_α .

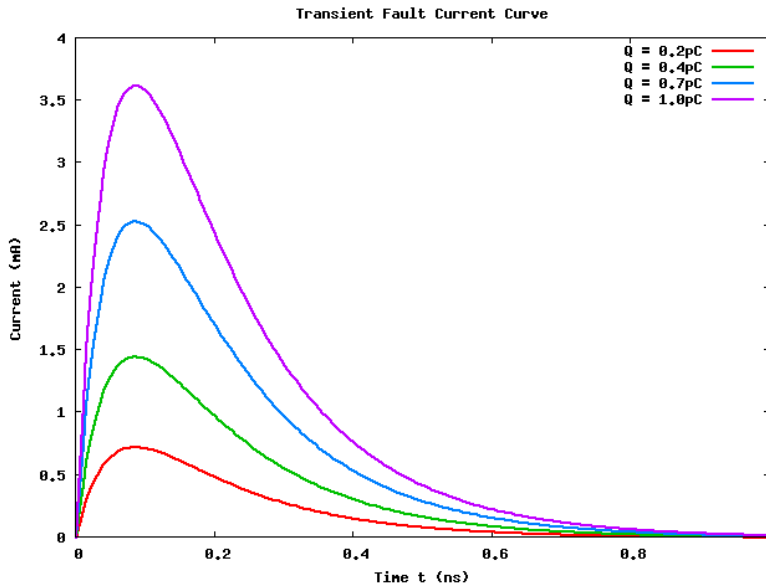


Figure 4.4: The current curve as result of a α -particle striking a device according [Mes82].

Figure 4.4 illustrates the current curve according the double exponential equation A.7 for $Q = 0.2pC$, $Q = 0.4pC$, $Q = 0.7pC$ and $Q = 1.0pC$. In these curves, τ_α and τ_β were defined as $1.06ns$ and $0.05ns$, respectively [DKCP94].

4.2.1.2 The Worst Case Depositing Charge

The amount of charge deposited by a particle in a device is widely discussed in the literature [ME02, GSB⁺04, Bau05]. These works consider the charged deposited by a

particle as a function of the energy of the ionizing particle and the device charge collection depth.

A particle with $1 \text{ MeVcm}^2/\text{mg}$ deposits approximately $10 \text{ fC}/\mu\text{m}$ of electron-hole pairs along each micron of its tracks. In addition, the LET of very few ionizing particles is higher than $15 \text{ MeVcm}^2/\text{mg}$ [ZM06].

A typical charge collection depth for heavy ions in a bulk silicon is approximately $2 \mu\text{m}$. Thus an ionizing particle with an LET of $15 \text{ MeVcm}^2/\text{mg}$ deposits approximately 0.3 pC of charge in any sensitive region it passes through.

Zhou and Mohanram present in [ZM06] a discussion about upper bounds for the deposited charges in some submicron process technologies. In other words, they suggest worst cases when choosing the worst case for a particle charge Q at ground level. For this, the linear energy transfer (LET) and the charging collection depth are taken into account. LET is a measure of the energy transferred to material as a function of an ionizing particle traversing through it.

Table 4.1: Summary of the worst case depositing charge for some technology processes for particles of with an LET of $15 \text{ MeVcm}^2/\text{mg}$

Technology process	Worst case depositing charge
180 nm	0.30 pC
130 nm	0.21 pC
100 nm	0.15 pC
70 nm	0.11 pC

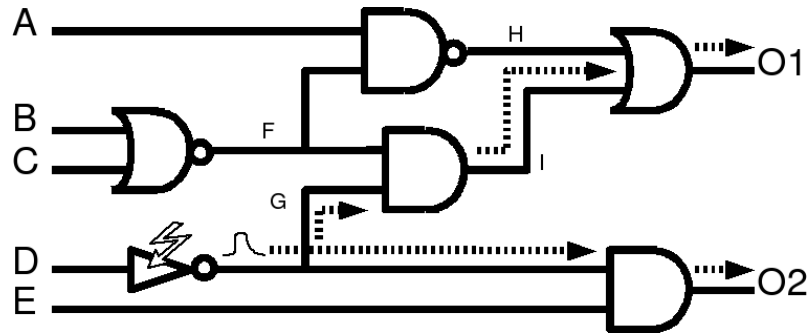
Table 4.1 summarizes the worst case depositing charge for some technology processes for particles with an LET of $15 \text{ MeVcm}^2/\text{mg}$ according [ZM06]. It is important to remark that the amount of deposited charge is not linearly reduced with the technology scaling.

They emphasize that in smaller technologies, the charge collection efficiency decreases due to the higher channel doping density and a reduction of the active layer thickness, which reduces the depletion width and channel funneling. Considering this characteristics of smaller feature sizes, their experiments give upper bounds of 0.21 pC , 0.15 pC and 0.11 pC for the 130nm , 100nm and 70nm process technologies.

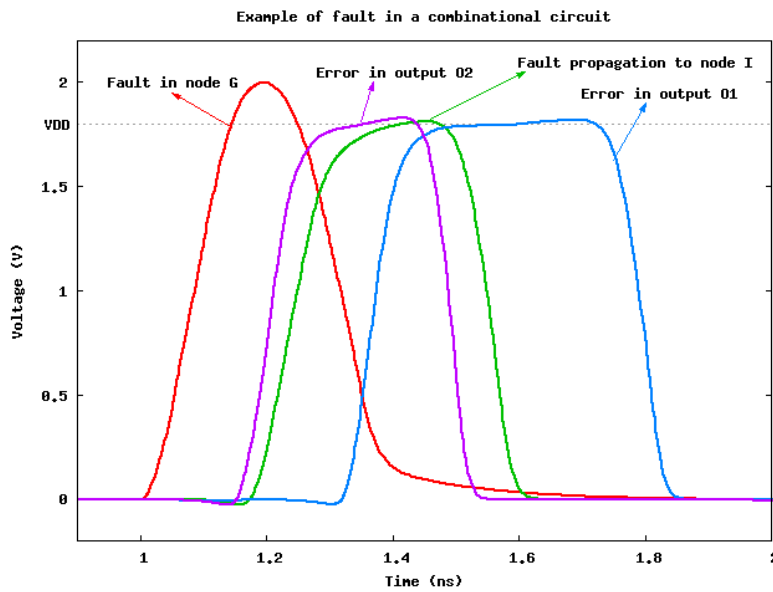
4.2.1.3 Single Event Transient (SET)

When a radiation-induced particle with enough energy hits a node of a circuit, a Single Event Transient (SET) pulse may be generated. Indeed, SETs are characterized as transient voltage fluctuations on circuit nodes. They can be caused by radiation-induced particles as previously discussed as well as by electrical noise in a noisy power supply, crosstalk noise, electromagnetic interference and radiation from lightning.

Figure 4.5 illustrates the effects of a particle striking a combinational block. Assuming the combinational block in Figure 4.5(a) with five inputs and two outputs and considering the logic values in the inputs of $A='1'$, $B='0'$, $C='0'$, $D='1'$ and $E='1'$. The Figure exemplifies an inverter being hit by a particle and the transient fault is propagated by the path



(a) Combinational block



(b) The fault propagation

Figure 4.5: The propagation of a transient fault as results of a particle striking a node of a combinational block.

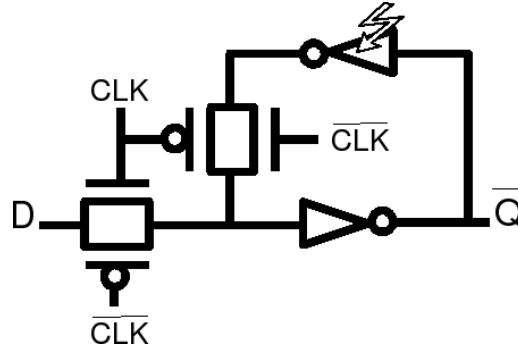
G and I . Thus, the outputs $O1$ and $O2$ have their states changed as consequence.

Waveforms are shown in Figure 4.5(b) where the transient pulse is propagated through the combinational circuit. The first path (node G to output $O1$ is longer than the second path (node G to output $O2$). Despite, the transient pulse is propagated through both paths provoking an error in the circuit outputs.

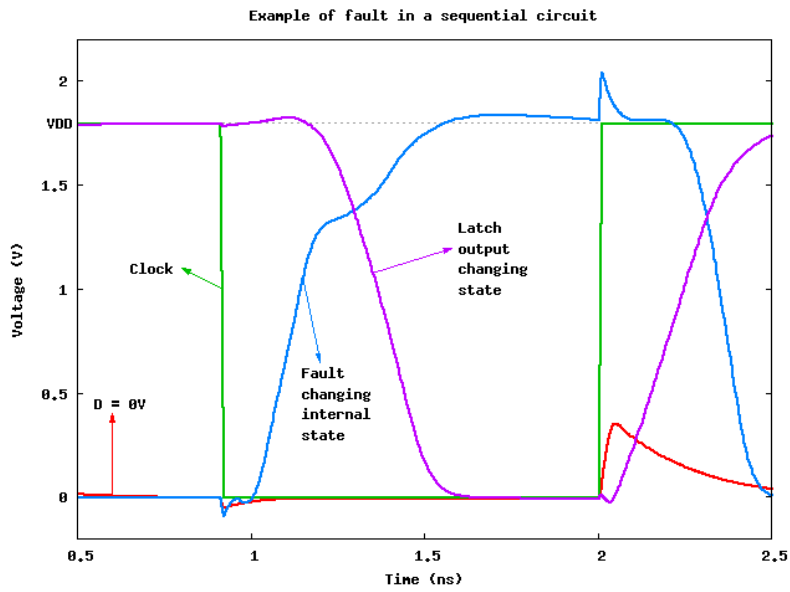
4.2.1.4 Single Event Upset (SEU)

A single event upset (SEU) is a change of state, or voltage pulse caused when a high-energy particle strikes a sensitive node in a micro-electronic device, such as in a micro-processor, semiconductor memory, or power transistors.

SEUs represent the radiation-induced hazard which is most difficult to avoid in space-borne applications, particularly in high density submicron CMOS ICs. Experimental results have shown that the critical charge collected at a sensitive node which is able to



(a) Classic latch



(b) The bit flip in the latch

Figure 4.6: Example of a bit flip in a classic latch.

produce an upset decreases as the inverse square of the feature size [NC97].

Figure 4.6 exemplifies a bit flip in a classic latch. The schematic of the latch is shown in Figure 4.6(a), which is composed by two inverters and two transmission gates controlled by the clock signal.

The waveform is shown in Figure 4.6(b). We consider the input D always set to “0” for this example. The single event disturbance starts at 1 ns at the internal node of the latch. At the same time, the output $\bar{Q}$ starts to change from V_{DD} to “0”. When the voltage in the internal node achieves $\frac{V_{DD}}{2}$, the inverters force the latch to change its state and the bit-flip occurs.

4.2.2 Single Hard Errors (SHE)

Single Hard Error (SHE) is as SEE, which causes a permanent change to the operation of a device. An example is a stuck bit in a memory device. In other words, we consider that a SHE occurs when the total dose of a single ion is sufficient to create a stuck bit.

Further details of the characterization of this kind of error in memories is presented in [DGC⁺92, PCBO94].

4.3 Other Radiation-Induced Effects

Single Event Latchup (SEL)

Single event Latchup (SEL) is a condition that causes loss of device functionality due to a single-event induced current state. During a SEL, the device exceeds the maximum current specified for the device. In this condition, it is mandatory that the power is removed. On the contrary, the device will be destroyed.

Single Event Burnout (SEB)

Single event burnout (SEB) is a condition that can cause permanent device destruction due to a high current state in a power transistor.

Single Event Gate Rupture (SEGR)

Single Event Gate Rupture (SEGR) is a single ion induced condition in power MOS-FETs which may result in the formation of a conducting path in the gate oxide.

4.4 Conclusion

This chapter presents basic aspects about radiation-induced effects in silicon devices. Some radiation effects are introduced in order to give some details about the behavior of a device with respect to a particle hitting it. The main radiation-induced effects are the total ionizing dose (TID), the displacement damage (DD) and the Single Event Effects (SEE).

The discussion is focused mainly in the Soft Single Event Effects due to their relevance with the development of this thesis. The main principles about Single Event Transients (SET) and Single Event Upsets (SEU) are highlighted.

The worst case depositing charge is also discussed, whose upper bounds for the critical charge as a function of some submicron technology processes are given. These values for the worst case depositing charge are used as basis to the development of the experiments in this work.

5 STATE-OF-THE-ART TECHNIQUES FOR SOFT ERROR PROTECTION

5.1 Introduction

Several techniques for soft error protection have been proposed in the literature in the last years. Most techniques are based on redundant structures. This redundancy may be temporal or spatial.

Temporal redundancy consists on inserting additional logic in the design, which guarantee the evaluation of a signal in different instants of operation. If an energetic particle hits a device creating a pulse voltage, this additional logic filters the signals and guarantee the attenuation of the pulse.

Spatial redundancy is usually based on the replication. The main idea in the spatial replication is that a particle hitting one of the elements does not affects the others. Thus, the outputs of the replicated elements are compared and the filtered signal is propagated.

All these techniques are characterized by the high area overhead and important delay and power penalties. The techniques are usually very effective against soft error effects, but the consequences are usually related to the inefficiency of the circuit concerning high frequency or power consumption.

Historically, sequential elements have been concerned for single event upsets. Efficient solutions to memory elements protection are presented in [BV94, CNV96, Roc88, WCL91]. However, since the transition time of the logic gates is getting shorter and clock frequencies are increasing significantly in nanometric technologies, errors in combinational logic parts are increasing and error rates will reach the same levels as in memories in the near future.

The constant decreasing size of microelectronics devices, associated with the reduction of voltage supply and the higher operating frequencies, leads to an increased vulnerability of logic circuits to soft errors [JJ05]. In [MT03], a study projects the soft error rate in combinational logic circuits comparable to unprotected memory elements by 2011.

This chapter presents some state-of-the-art techniques to mitigate the effects of single events in sequential and combinational circuits. The advantages and drawbacks of each technique are also discussed. Besides, a temporal redundancy method is proposed to protect sequential elements against SEEs.

5.2 The Classic Techniques

Many techniques have been presented in the literature aiming at protecting combinational blocks and storage elements against Single Event Error (SEE). The techniques are usually based on redundancy of elements or sizing transistors.

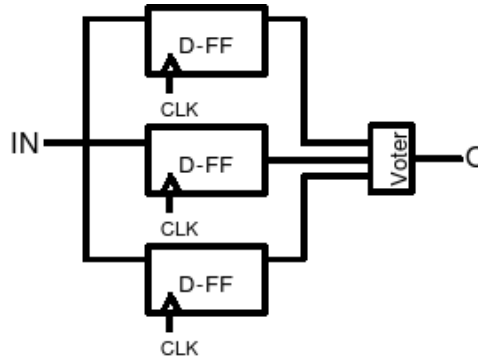


Figure 5.1: The classic TMR.

Triple Modular Redundancy (TMR) is the most common used technique to protect sequential circuits. The TMR technique consists on triplicating elements (spatial redundancy) in such way the logic value resulting from at least two elements are propagated. Figure 5.1 shows the classic TMR technique applied to D-FlipFlops. This technique can be used to prevent an error as a result of a single fault occurring inside the elements. On the other hand, when a SET occurs before the inputs (e.g. a combinational block), the transient pulse may be captured by the flipflops.

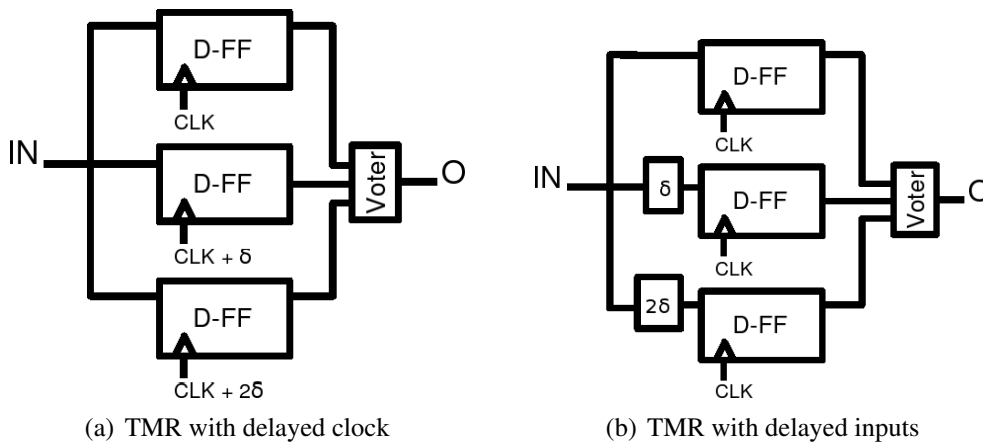


Figure 5.2: Two TMR versions with delayed clock and delayed inputs to avoid transient faults coming from the combinational blocks.

A way to provide perturbation tolerance on both combinational and sequential blocks is using the TMR technique in such way we can guarantee that the fault is only propagated

to one Flip-Flop. In other words, providing perturbation tolerance by inserting temporal redundancy.

The Figure 5.2 shows two implementations of the TMR technique targeting perturbation tolerance to combinational and sequential blocks. In the first implementation (Figure 5.2(a)), the clock signal of each Flipflop (CLK , $CLK + \delta$ and $CLK + 2 \times \delta$) is delayed in such a way the input signal is captured at three different moments. Thus, a transient fault at the input is captured only by one of the flipflops.

The same idea is used in Figure 5.2(b), where the same clock signal is used in the tree Flip-Flops, but the input signal is delayed by *Delay Blocks*. The time penalty of these TMR techniques over the time of a D-Flipflop is $2 \times \delta + T_{DelayVoter}$.

The TMR is usually applied to sequential elements but it can be used to combinational blocks and a whole circuit but they present a very high overhead in area (more than 200%). Furthermore, techniques that modify clock signal as shown in Figure 5.2(a) may present additional and usually unnecessary design challenges due to clock skew and clock tree distribution.

5.3 Gate Duplication Methods

Many techniques based on gate duplication have been proposed to reduce the soft error failure rate in combinational logic. The gate duplication consists of connecting inputs and outputs of a gate with an identical gate.

A partial gate duplication method for soft error failure rate reduction in combinational circuits is presented in [MT03]. The partial gate duplication consists of selectively duplicating the most sensitive gates.

Results show that a 90% soft error failure rate reduction for energetic particles of 10 MeV is obtained with 50% area overhead. The drawback of this method is related to the area overhead to achieve a 0% sensitivity circuit.

The area overhead in this methodology presents an exponential behavior, which small area overhead is obtained with up to 80% soft error reduction rate. Soft error reduction rates higher than 80% result in area overheads close to 100%.

A SER analysis in combinational logic circuits and a partial gate duplication methodology for soft error protection is proposed in [HN06, NJJ06]. These works present a very important contribution in relation to SET analysis. The SET propagation is analyzed with respect to logical and electrical masking.

They consider that the logical and electrical masking analysis contribute in the SER of the circuit. Thus, the logical masking contributes by given the probability of a single event to be masked by the circuit logic and the electrical masking contributes to the circuit SER by analyzing the electrical attenuation of the single event. More details about the logical and electrical masking are given in the Chapter A.4.

The drawback of this technique is mainly related to the resulting overhead. The duplication of a gate increases twice the occupied area, but also increases the input capacitances. These capacitances increase the delay of the connected gates. This factor causes disturbs in the delay paths of the circuit and makes it difficult the timing closure.

5.4 Gate Sizing Techniques

The gate sizing method consists on changing the cells of the circuit with bigger ones with the same logic function. A library of cells is composed by some versions of each gate. Thus, a sensitive cell may be changed by a more robust one in order to reduce the sensitivity.

Gate sizing techniques for soft error failure rate are presented in [DDCS04, CRO⁺05, ZM06]. The contribution of these works are related to the techniques to analyze the sensitivity of the circuits and the selection of the gates. The sizing of the gates is done with analytical model, which are computationally efficient.

In [ZM06], candidate gates are selected according to the sensitivity level. The sensitive level is determined by the logical masking of each node in the circuit.

Design penalties of the gate sizing technique are positives in comparison with the duplication method due to the reduction on the area overhead. However, the reduction of penalties could be more efficient if the sizing algorithm could size pull-up and pull-down blocks of the gates.

A new sizing algorithm is proposed in the Chapter A.4. This algorithm takes into account different characteristics of PMOS and NMOS transistors in order to reduce the penalties provoked by a radiation-hardened design.

5.5 Protecting Sequential Elements Through Feedback Control

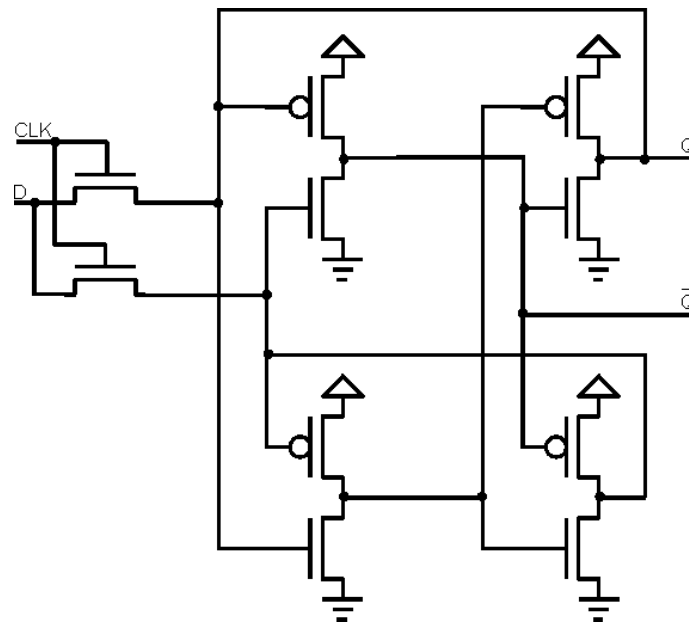


Figure 5.3: Latch using the DICE technique as proposed in [CNV96].

An interesting technique to protect storage elements is the Dual Interlocked storage Cell (DICE) [CNV96]. The main idea of this technique relies on the principle of *dual node feedback control* in order to achieve immunity to upsets. In other words, this means

that the logic of each node of the cell is controlled by two nodes. Thus, a perturbation in a node is removed after the upset due to the state-reinforcing feedback ensured by the other two nodes.

Figure 5.3 shows a latch using the DICE technique, where each inverter will change its state only if two other inverters change their state. This technique presents an important characteristic for the design of SRAM memories with upset immunity. Further details can be seen in [CNV96].

5.6 Using Time Redundancy To Protect Sequential Elements

In [Ang00], it is presented a technique aiming at taking advantage from the temporal nature of transient faults, and achieve transient fault tolerance by using time redundancy. This technique should lead to a significant reduction of hardware cost compared to the TMR classic solution because the main idea is to combine self-checking design with time redundancy.

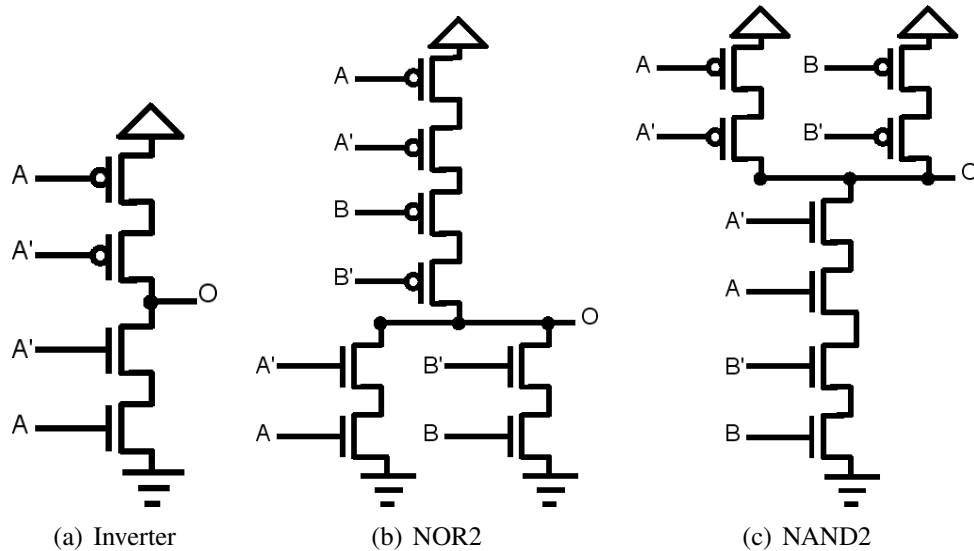


Figure 5.4: INV, NOR2 and NAND2 gates using the CWSP technique proposed in [Ang00].

The fact that soft-errors affect the outputs of the circuit only for short time duration can be exploited by using asynchronous sequential elements. These elements produce on the outputs a determined state for each correct input. This state corresponds to the circuit fault-free operation. In addition, the element preserves its previous output state for each erroneous input. In addition, if a transient pulse changes a fault free input into an erroneous one, the output state produced by the correct input is preserved.

A way to generate a *Code Word State Preserving* (CWSP) element is to replace each transistor of the gate by a pair of transistors connected in series and driven by duplicated inputs. In this gate, when the inputs of a pair of transistors are equal, the two transistors behave as a single transistor driven by one of the duplicated inputs.

When the inputs of one or more transistor pairs have not equal values due to a transient error, the two transistors of the pairs behave as a single transistor in off state. This situation preserves the same state at the output of the CWSP element, as the same state obtained before the transient errors drive some input into non-equal values.

Figure A.9 shows examples of the inverter, NOR and NAND gates using this principle. In the time redundancy approach, instead of circuit duplication we can duplicate the output signal of the circuit in the time domain, by observing this signal at two different instants. One of the inputs of the CWSP element is coming directly from the combinational circuit output while the other input is delayed.

Table 5.1: Area and delay of the library cells INV, NAND and NOR in comparison with CWSP cells automatically generated with the tool presented in [LDGR03].

	Area (μm^2)			Delay (ps)		
	Typ.	CWSP	Over.	Typ.	CWSP	Over.
INV	8.19	11.64	42%	83	160	92%
NAND	12.28	19.07	55%	98	172	75%
NOR	12.28	19.07	55%	102	260	154%

Table A.3 presents the occupied area and propagation time of CWSP cells according to Figure A.9 in comparison of typical cells presented in a $0.18\mu m$ standard cell library. It is shown that the area overhead is between 42% and 55% and the propagation time is between 92% and 154% applying the same output capacitance to both gates. The tool presented in [LDGR03, SWL⁺03, BLGR04] was used to automatically generate these cswp cells.

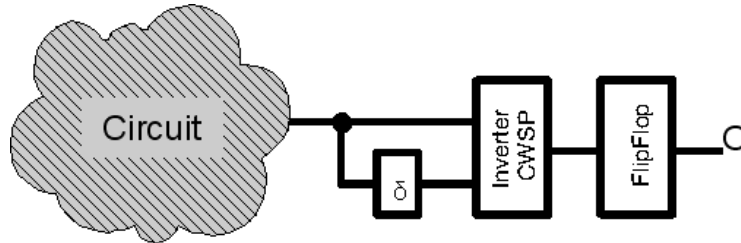


Figure 5.5: Perturbation tolerant circuit based on time redundancy.

The implementation of perturbation tolerant combinational circuits based on time redundancy is presented in Figure A.10. The delay block must be able to degrade the signal at the input of the CWSP cell according to the time of the transient fault that we achieve to tolerate. The time penalty in this case is $D_{cw} + 2 \times D_{tr}$, where D_{cw} is the logic transition time and D_{tr} is the duration of the transient pulse.

Table A.4 presents the total area and propagation time of the new CWSP cells developed as shown in Figure A.10. The development of CWSP cells supporting transient faults ranging from $250ps$ to $500ps$ are assumed. Thus, delay blocks were developed and inserted in the CWSP gates in order to obtain hardened cells.

Table 5.2: Total area and delay of CWSP cells.

	Trans. (ps)	Area (μm^2)	Delay (ps)
INV	250	28.8	323
	500	46.0	538
NAND	250	59.2	370
	500	91.2	559
NOR	250	59.2	352
	500	91.2	572

We proposed in [LAR05a] a technique targeting perturbation tolerance to both combinational and sequential logics. It uses the timing redundancy technique presented in [Ang00] to provide fault tolerance in circuits. To tolerate transient faults in combinational circuits as well as in the sequential elements (e.g. latches) our approach uses a modified latch where the last inverter stage of the combinational circuit is replaced by a CWSP inverter, while a delay block has been introduced in the feedback path of the latch.

Figure A.11 shows the implementation of a latch using the CWSP logic. The technique uses a CWSP inverter and *Delay Blocks* intend to achieve fault tolerance through timing redundancy. It was proved by transient fault simulation that these CWSP d-latches have 100% transient fault coverage, assuming that the *delay blocks* have a delay greater than the time of the transient pulse.

Table 5.3: Area overhead comparison of TMR and CWSP d-flipflops.

	Area	
	μm^2	Overhead
Standard Cell	57.6	—
CWSP (250ps)	181.7	215%
CWSP (500ps)	249.6	333%
TMR (250ps)	206.1	258%
TMR (500ps)	206.1	258%

Table A.5 presents the comparison between a classic D-FlipFlop found in a $0.18\mu m$ standard cell library and the transient robust latch proposed in [LAR05a] (Figure A.11) and the TMR FlipFlop. We assume that delay blocks in the TMR D-FlipFlops can be shared between all FlipFlops in the same clock domain, reducing the occupied area. Despite of that, results show that the robust CWSP FlipFlops present smaller area overhead against faults of 250ps in comparison to the TMR technique.

A case study was done in order to verify the penalties of the insertion of CWSP cells in a MIPS-like processor and a 8051 controller. The logic synthesis and technology mapping were done with Synopsys Design Compiler and the whole circuit layout was generated with Cadence Silicon Ensemble. The CWSP robust latch layout was generated by using Parrot Punch tool and inserted in the system layout.

Table A.6 shows the occupied area and frequency in the MIPS and 8051 architectures

Table 5.4: CWSP and TMR techniques on microprocessors.

	MIPS				8051			
No. Comb. El.	11,968				5,408			
No. Flipflops	1,793				1,359			
	Area		Frequency		Area		Frequency	
	μm^2	Over	<i>MHz</i>	Penalty	μm^2	Over	<i>MHz</i>	Penalty
Classic	480,317	—	77.7	—	234,720	—	58.2	—
CWSP (250ps)	746,480	55%	75.8	2.4%	436,240	85%	57.2	1.8%
CWSP (500ps)	890,172	85%	72.7	6.7%	550,560	134%	55.4	5.0%
TMR (250ps)	808,200	68%	73.9	5.1%	491,400	109%	56.0	3.8%
TMR (500ps)	808,200	68%	71.2	9.0%	491,400	109%	54.5	6.8%

flipflops by inserting CWSP elements in their structures. A case study was done in order to validate this technique and it is presented in Section 5.6.

Table 5.5 presents an overview of the techniques presented in this chapter, where advantages and drawbacks are highlighted.

Table 5.5: Review of the techniques presented in this chapter.

Technique	Description, advantages and drawbacks
Classic TMR	Consists on triplicating elements. A voter propagates the state coming from two or more replicated elements. TMR works whenever the single event occurs in only one element, independently of the particle energy. Faults may be captured by the inputs. Timing penalties are insignificant but it presents high area and power overhead, mainly if the whole circuit is triplicated. (Figure 5.1)
TMR w/ delayed clock	Consists on triplicating elements with the delayed clocks $CLK, CLK + \delta$ and $CLK + 2 \times \delta$. Transient pulses with duration smaller than δ are captured by only one element. Additional complexity is inserted in the design due to the clock management. (Figure 5.2(a))
TMR w/ delayed inputs	Consists on triplicating elements with delay blocks in the inputs. Transient pulses with duration smaller than δ are captured by only one element. The technique does not insert additional complexity in the clock synthesis but increases the clock period by $2 \times \delta$. (Figure 5.2(b))
Gate Duplication	Gate duplication places two gates, whose inputs and outputs are connected. The gates are selectively duplicated according to sensitivity levels. The main drawback is related to area overhead and delay penalties. (Section 5.3)
Gate Sizing	The gate sizing technique consists on changing the cells by bigger ones. The area overhead is less significant than gate duplication but delay penalties are also important. (Section 5.4)
Feedback Control	Feedback control relies on the principle that dual nodes together are able to attenuate the single event. The area overhead is considerably smaller than the TMR and time penalties are insignificant. Faults may be captured by the inputs. (Section 5.5)
CWSP	Used in combinational blocks. The CWSP techniques replaces each transistor of the gate by a pair of transistors connected in series and driven by duplicated inputs. A block with delay δ must guarantee that a transient pulse arrives in only one of the inputs. Area overhead is insignificant, but clock period increases by δ (Section 5.6)
CWSP Sequential	This technique inserts CWSP elements inside sequential structures as flipflops and latches. The sequential element is able to filter single events with duration smaller than δ at internal blocks and coming from the combinational block. The area is increased by the insertion of two delay blocks, and the clock period increases by δ . (Section 5.6)

6 AN EFFICIENT TRANSISTOR SIZING METHODOLOGY FOR SOFT ERROR PROTECTION IN COMBINATIONAL LOGIC CIRCUITS

6.1 Introduction

A radiation-induced effects overview is given in Chapter 4. The chapter introduces the causes and effects of the radiation-induced effects and shows a way to model the behavior of current injected by a particle in the devices.

State-of-the-art methods for soft error failure rate reduction are presented in Chapter A.3. These techniques are very effective on protecting circuits against single event upsets. However, radiation-hardened designs are usually inefficient concerning area, delay or power consumption. Most techniques fail in these three aspects.

A new transistor sizing method is presented in this chapter. The main characteristic of the proposed methodology is to find the smallest transistors width to attenuate SETs in the nodes of a combinational circuit.

Another important point is that pull-up and pull-down transistors are independently sized, minimizing the area overhead and the power consumption. In other words, we apply asymmetric transistor sizing to attenuate SETs with minimized area overhead. Works presented in the literature are based in symmetric models to size pull-up and pull-down blocks.

6.2 Combinational Circuits Sensitivity

The sensitivity analysis of circuits has been presented in several works [NJJ06]. Most of them include the structure of the gates and layout details. The analysis of the structure of a gate consists on evaluating the propagation of a fault as a functions of transistors connections. For example, a fault in the drain node of a transistor has more probability to be propagated than a node connected to the V_{DD} or GND .

The sensitivity analysis considering the layout takes into account the probability of a particle to hit a region of the layout. For example, a big drain area has as higher probability to be hit than a smaller one.

In this work, the structure of the gate and its layout is not considered. Differently, we consider only the output node of each cell due to its higher sensitivity in comparison with the internal nodes of the gate. Layout characteristics are not taken into account in

Table 6.1: Probability of a node as a function of the gate equation [JJ05].

Logic Function	Resulting Probability
AND	$P_Z(1) = P_a(1) * P_b(1)$
NAND	$P_Z(1) = 1 - P_a(1) * P_b(1)$
OR	$P_Z(1) = 1 - (1 - P_a(1)) * (1 - P_b(1))$
NOR	$P_Z(1) = (1 - P_a(1)) * (1 - P_b(1))$
XOR	$P_Z(1) = P_a(1) + P_b(1) - 2 * P_a(1) * P_b(1)$
XNOR	$P_Z(1) = 1 - P_a(1) - P_b(1) + 2 * P_a(1) * P_b(1)$
BUF	$P_Z(1) = P_a(1)$
INV	$P_Z(1) = 1 - P_a(1)$

the sensitivity analysis because we consider the sensitivity of a gate after sizing becomes zero as function of a given critical charge Q_c .

It is important to remark that layout aspects are not considered only for sensitivity analysis purposes. Layout details are essential for the proposed transistor sizing methodology.

We consider the logical and electrical masking as the sensitivity of a circuit. The logical masking represents the probability of a transient pulse to be masked by the logic function of the circuit, and the electrical masking describes if a transient pulse in a node is not propagated to the primary outputs (PO) or flipflops. Thus, the sensitivity of a circuit is given by

$$S_{circuit} = \sum_{n=1}^N (1 - L_n) \cdot (1 - E_n) \quad (6.1)$$

where L_n is the logical masking and E_n is the electrical masking. The logical masking L_n is a probability value. Bigger logical masking means smaller probability of a transient pulse to be detected in the circuit outputs. The electrical masking E_n is a boolean value where “0” means that the transient pulse is totally attenuated and “1” indicates that the transient can be seen in the outputs.

6.2.1 The Logical Masking

The *logical masking* occurs when a SET provoked by a particle is not propagated to a primary output (PO) due to the logic of the circuit. In other words, the pulse is masked as function of the vector applied in the primary inputs (PI) of the circuit. Controllability and observability techniques are used to define the logical masking of a node.

Controllability in combinational logic circuits denotes the ability to a state be set in a node. Observability is a measure of how well a state in a internal node can be known at the primary outputs (PO).

The controllability of the gate output node is obtained by the logic function of the gates as shown in Table A.7 [JJ05]. Thus, the propagation of the controllability probability is done for the entire circuit, from the PIs through each gate until the POs are reached.

Figure A.12 illustrates the logical masking in a gate. A pulse in one of the gate inputs is propagated through the gate only if a non-controlling value is applied at the other input.

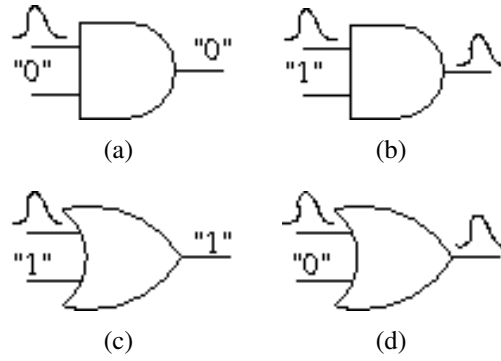


Figure 6.1: The logical masking.

Figure A.12(a) shows the logical masking in the AND gate as a function of the controlling logic value “0” at the input. Otherwise, the logical masking does not happen if a non-controlling value is applied (Figure A.12(b)).

In an OR gate, the same situation is considered, where the pulse propagates through the gate only if the non-controlling value is applied to the other input. Figure A.12(c) shows the logical masking as function of a controlling value and Figure A.12(d) shows the case where there is no logical masking.

6.2.2 The Electrical Masking

Electrical masking can be defined as the electrical attenuation of a pulse in a node by the gates in a path to the point that the SET does not affect the result of a circuit.

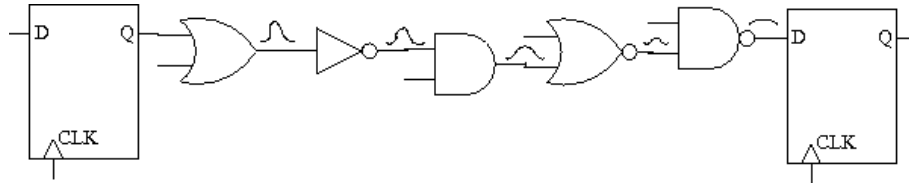


Figure 6.2: The electrical masking.

Figure A.13 shows an example of SET degradation. This degradation is the base of the electrical masking, where the pulse is degraded as a function of the electric characteristics of the gates in the path. The pulse can be captured by the memory element if it is not enough degraded. More details about electrical masking and SET propagation are given in Section A.4.2.3.

6.3 An Analytical Model for Single Event Transients

The sensitivity model used in our transistor sizing strategy was proposed in [WVK07]. The model is based on two electrical device parameters. The effective loading capacitance C lumped onto the output node of a gate g and the effective resistance R of the “ON” transistors of this gate.

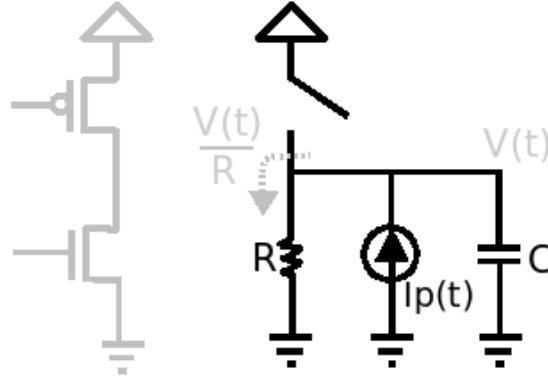


Figure 6.3: Equivalent circuit for calculating circuit response to an energetic particle hit.

For modelling purposes, circuit response for the energy particle is modeled as the network depicted in Figure A.14, and may be represented by

$$\frac{V(t)}{R} - I_p(t) - C \frac{dV(t)}{dt} = 0 \quad (6.2)$$

where the term $\frac{V(t)}{R}$ consists of the discharging current in the transistor, represented by the resistor R . $I_p(t)$ represents the current caused by the particle hitting the device and the last term $C \frac{dV(t)}{dt}$ represents the current in the capacitor C .

The model derivation has a strong relation with the electrical devices behavior and allows the evaluation of the critical charge Q_c needed to induce a SET in a node, and the transient pulse duration, as well.

6.3.1 Modeling Resistances and Capacitances

The use of linear resistors to model transistor paths is a widely known method [WE93]. Thus, the effective resistance R can be analytically determined by

$$R = \frac{1}{\mu_0 C_{ox} \left(\frac{W}{L} \right) (V_{gs} - V_{th})} \quad (6.3)$$

where μ_0 is the mobility of the transistor channel. C_{ox} is the oxide capacitance, which is given by $\frac{\epsilon_0 \epsilon_{SiO_2}}{t_{ox}}$. ϵ_0 is the dielectric constant, ϵ_{SiO_2} is the oxide relative dielectric constant and t_{ox} is the gate oxide thickness. V_{gs} is the gate-source voltage and V_{th} is the threshold voltage.

All these parameters are constants related with the technology process, except by the $\left(\frac{W}{L} \right)$ ratio that represents the transistor dimensions. Based on this aspect ratio, we are able to explain the relation between the transistor width and the resistance. The smaller the transistor width, the higher the resistance.

Figure A.15 illustrates two stacked transistors modeled as resistors. Assuming NMOS transistors are “ON” in the NAND gate of the example due to input signals $a = “1”$ and $b = “1”$. The effective resistance R is given by the sum of the resistances r_1 and r_2 .

The effective capacitances C is defined by the sum of three capacitances connected to the output node.

$$C = C_{diffusion_{g1}} + C_{connection} + C_{gate_{g2}} \quad (6.4)$$

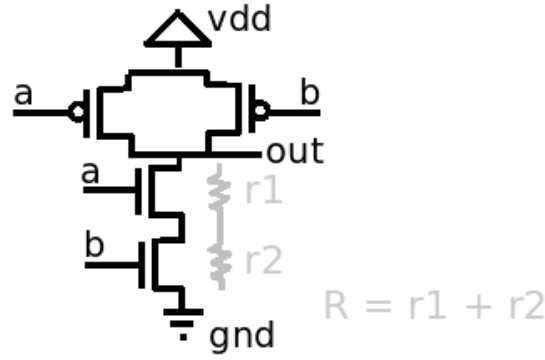


Figure 6.4: A transistor modeled as a resistance.

Table 6.2: Approximation of intrinsic MOS gate capacitance.

Parameter	Off	Non-saturated	Saturated
C_{gb}	$C_{ox}A$	0	0
C_{gs}	0	$\frac{1}{2}C_{ox}A$	$\frac{2}{3}C_{ox}A$
C_{gd}	$C_{ox}A$	$\frac{1}{2}C_{ox}A$	0
$C_g = C_{gb} + C_{gs} + C_{gd}$	$C_{ox}A$	$C_{ox}A$	$\frac{2}{3}C_{ox}A$

where $C_{diffusion_{g1}}$ is the sum of all PN junction capacitances of the driving gate. $C_{connection}$ is the wiring and parasitic capacitances, and $C_{gate_{g2}}$ is the gate capacitance of all transistors connected to the output node.

The $C_{diffusion_{g1}}$ is given by

$$C_{diffusion_{g1}} = \sum_d^D \times C_{ja}A_d + C_{jp} \times P_d \quad (6.5)$$

where C_{ja} is the junction capacitance per μ^2 , A_d is the diffusion area, C_{jp} is the periphery capacitance per μ and P_d is the diffusion perimeter.

The third term of the capacitance C is the gate capacitance. Thus, $C_{gate_{g2}}$ is defined according to the region the gate $g2$ is operating.

Table A.8 presents the gate capacitance according to the region of operation. Based on this information, the gate capacitance is defined by

$$C_{gate_{g2}} = \sum_{g_{off}} C_{ox}A_g + \sum_{g_{on}} \frac{2}{3}C_{ox}A_g \quad (6.6)$$

These analytical equations allow to model the behavior of the transient pulse as a function of the electric characteristics of the devices.

6.3.2 The Single Event Transients Model

The single event model uses the double exponential equation discussed in Section 4.2.1.1. Important characteristics about the transient pulse can be obtained by (A.7).

Models presented in [WVK07] are derivations of the double exponential to obtain the peak time t_{peak} and the voltage peak V_{peak} .

It is important to remark that the τ_β is considered to be much smaller than τ_α ($\tau_\alpha \gg \tau_\beta$) in the formulations. In other words, the model assumes a very fast rise time to the double exponential.

This differential equation (A.2) can be solved in order to obtain the voltage $V(t)$ at the struck node. Thus, $V(t)$ is given by

$$V(t) = \frac{I_0 \tau_\alpha R}{\tau_\alpha - RC} (e^{\frac{-t}{\tau_\alpha}} - e^{\frac{-t}{RC}}) \quad (6.7)$$

The time t_{peak} at which the node voltage reaches its maximum value can be evaluated by

$$t_{peak} = \frac{\ln\left(\frac{\tau_\alpha}{RC}\right) \tau_\alpha RC}{\tau_\alpha - RC} \quad (6.8)$$

and, inserting (6.7) into (A.8) leads to the peak transient voltage V_{peak} reached at the struck node.

$$V_{peak} = \frac{I_0 \tau_\alpha R}{\tau_\alpha - RC} \left(\left(\frac{\tau_\alpha}{RC} \right)^{\frac{RC}{RC - \tau_\alpha}} - \left(\frac{\tau_\alpha}{RC} \right)^{\frac{\tau_\alpha}{RC - \tau_\alpha}} \right) \quad (6.9)$$

where R is the effective resistance of the pull-up path (if PMOS transistors are “ON”) or the effective resistance of the pull-down path (if NMOS transistors are “ON”) and C is the effective capacitance loading lumped onto the output node.

The critical charge Q_c can be derived from (A.9) once the V_{peak} of a transient pulse is known. Thus, the critical charge Q_c is given by

$$Q_c = \frac{V_{peak}(\tau_\alpha - RC)}{R \left(\left(\frac{\tau_\alpha}{RC} \right)^{\frac{RC}{RC - \tau_\alpha}} - \left(\frac{\tau_\alpha}{RC} \right)^{\frac{\tau_\alpha}{RC - \tau_\alpha}} \right)} \quad (6.10)$$

The voltage at the struck node shows a double exponential behavior in which the transient voltage V_{peak} is reached at time t_{peak} . The voltage starts to decrease exponentially after t_{peak} .

$$\tau_n = t_{peak} - RC \ln \left(\frac{\frac{1}{2} VDD}{V_{peak}} \right) - \tau_\alpha \ln \left(\frac{\frac{1}{2} VDD}{V_{peak}} \right) \quad (6.11)$$

Equation (A.11) shows the transient pulse duration τ_n , where the second term corresponds to the analytical solution if RC time is much greater than τ_α and the last term corresponds to the analytical solution if τ_α time is much greater than RC .

6.3.3 Single Event Transient Propagation

The analysis of the transient pulse propagation shows that the pulse degradation is directly influenced by the propagation delay τ_g . In other words, larger τ_g leads to greater degradation of the transient pulse.

Wirth *et al* proposed a pulse degradation model based on curve fitting [WVKNK07]. The model considers a k parameter equal to the minimum ratio τ_n/τ_g needed to propagate a SET to the next stage in a circuit path.

The transient pulse degradation is given by

$$\tau_{n+1} = \begin{cases} 0 & \text{if } (\tau_n \leq k\tau_g), \\ (k+1)\tau_g(1 - e^{-(\tau_n/\tau_g)}) & \text{if } ((k+1)\tau_g < \tau_n \leq (k+3)\tau_g), \\ \frac{\tau_n^2 - \tau_g^2}{\tau_n} & \text{if } ((k+1)\tau_g < \tau_n \leq (k+3)\tau_g), \\ \tau_n & \text{if } (\tau_n > (k+3)\tau_g). \end{cases} \quad (6.12)$$

For an input transient with small duration, the output voltage peak does not reach $\frac{1}{2}VDD$ and complete attenuation must be considered. Thus, the first case models situations where the transient pulse is totally suppressed.

Second and third cases are related to a partial degradation in the transient pulse according to the relation with the SET duration τ_n and the gate delay τ_g . The fourth degradation case consists on situations where the pulse is not degraded from a stage to another or it can be neglected.

These four degradations cases are the basis to the sizing algorithm because of its propagation properties. These properties can be useful also to obtain the maximum acceptable transient pulse duration in a node.

We consider maximum acceptable transient pulse in a node the maximum SET duration that is attenuated before the primary outputs. In other words, the maximum acceptable SET duration is a pulse in which is not propagated to any PO.

One important remark is that a transient pulse does not need to be attenuated in the net n (except in cases where gates are connected to outputs). The SET may be attenuated through the gates in the whole path between the node n and the primary outputs. The complete attenuation of a SET in a net results in unnecessary oversized transistors.

6.4 The Transistor Sizing Strategy

The transistor sizing strategy proposed in this thesis consists of finding the smallest transistor width of each circuit gate for SET attenuation. The formulations previously discussed are the basis to the sizing algorithm.

The proposed transistor sizing strategy is presented in Algorithm 9. First lines (2-6) define the circuit sensitivity as shown in (A.1). The transistor sizing strategy starts at line 8, where every node n of the circuit is visited in order to find the minimum transistor width to each gate g connected to this node. It is important to note that only nodes with the sensitivity bigger than the maximum defined sensitivity M are sized (line 11).

Function `getMaximumSET(  $n$ ,  $g$  )` (line 12) finds the maximum pulse duration τ_n in the node n that is suppressed before the primary outputs. The transistor sizing algorithm to a gate g is function of this SET duration τ_n .

Function `sizeTransistors(  $s$ ,  $g$ ,  $\tau_n$  )` (line 13) continuously increases the transistors width until the SET in the node n be smaller than τ_n . When this situation is reached, we consider the transistors of the gate g are sized as expected to the charge Q_c .

The implementation of these functions are discussed in details in Section A.4.4.

Other lines of the strategy shown in Algorithm 9 give some idea about the navigation in the nets. The algorithm evaluates every node of the combinational logic, from the

Algorithm 8 The transistor sizing for SET attenuation.

Require: Set of gates G , Set of Nets N , Set of outputs O , Maximum sensitivity M , Max critical charge Q_c , Desired circuit sensitivity $S_{desired}$

Ensure: Set of gates with sized transistors G_{new}

```

1:  $G_{new} \leftarrow \emptyset$ 
2: for all  $n \in N$  do
3:    $L_n \leftarrow \text{calculateLogicalMasking}(n)$ ;
4:    $E_n \leftarrow \text{calculateElectricalMasking}(n, Q_c)$ ;
5:    $S_n \leftarrow (1 - L_n) \cdot (1 - E_n)$ 
6: end for
7:  $V \leftarrow O$     {Nets to visit, starting from the outputs.}
8: while  $V \neq \emptyset$  do
9:   for all  $n \in V$  do
10:     $g \leftarrow \text{getFaninGateConnectedToNet}(n)$ ;
11:    if  $S_n > M$  then
12:       $\tau_n \leftarrow \text{getMaximumSET}(n, g)$ ;
13:       $g_{new} \leftarrow \text{sizeTransistors}(s, g, \tau_n)$ ;
14:       $G_{new} \leftarrow G \cup \{g_{new}\} \setminus \{g\}$ 
15:    end if
16:     $I \leftarrow \text{getGateInputs}(g)$ ;
17:     $V \leftarrow V \cup I \setminus \{n\}$ 
18:   end for
19: end while

```

primary outputs (PO) to the primary inputs (PI). This is done because the delay of the gates is changed after sizing. When transistors of a gate are sized, the delay usually becomes smaller and a transient pulse propagates with smaller degradation.

Erroneous interpretation concerning the SET propagation must happen if the transient pulse is evaluated before the sizing of the gates in the path to the POs. Thus, when the SET is evaluated in a node n , we guarantee that every gate in the path between this node n and the POs were already sized.

6.5 The Transistor Sizing Model

The transistor sizing technique proposed in this thesis is basically separated in three steps. These steps are related to the sensitivity of a circuit node as a function of a particle hitting the circuit and the minimum transistor width needed to attenuate a SET (electrical masking).

The first step is the sensitivity analysis, which is represented in lines 2 to 6 (Algorithm 9). The sensitivity of a node n is given by the logical and electrical masking.

Assume the example in Figure A.16 where a particle with charge Q hits the output of the NAND gate at the *stage 0*. The transient pulse is propagated from the net b through each stage to the circuit output. The transient pulse duration at the struck node is defined by equations described in Section A.4.2.2 and is a function of the resistance R , the capacitance C , the charge Q and some technology process constants.

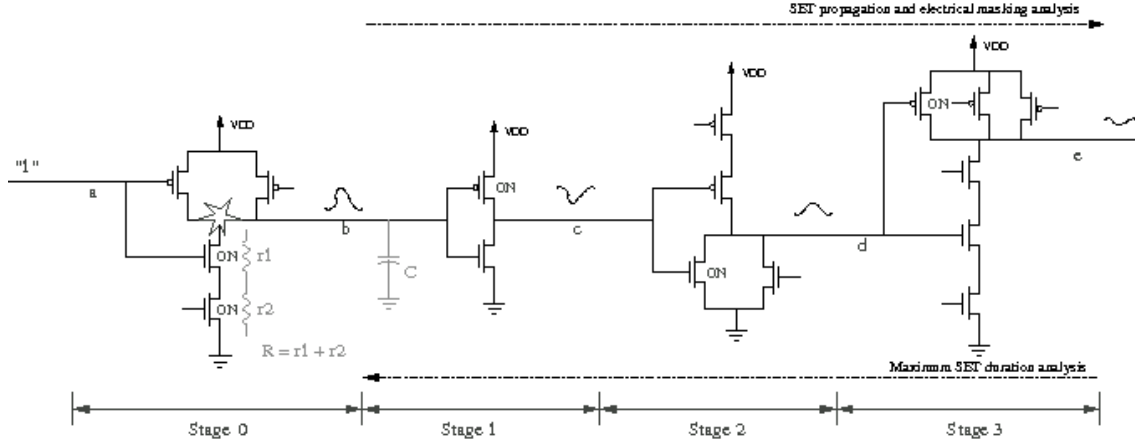


Figure 6.5: A transient pulse propagation example.

The degradation of a SET depends on the delay of each gate in the path as explained in Section A.4.2.3. The electrical masking is basically the analysis of the SET degradation through the path. A node is assumed to be electrically masked if the transient pulse is suppressed before the outputs.

For this reason, function `getMaximumSET( n, g )` (line 18) finds the maximum SET duration τ_n for a net n that is attenuated just before the primary outputs. Assuming SET duration at the primary outputs $\tau_{out} = 0$, equations from Section A.4.2.3 were modified to find an acceptable SET duration in the net n . In other words, we derived those equations to obtain the SET duration τ_n at the inputs of each gate as a function of the SET duration τ_{n+1} at its output.

The transient pulse propagation equations in Section A.4.2.3 consider four situations where the SET duration may be totally degraded, partially degraded or propagated.

Cases where the pulse is totally attenuated or propagated without any degradation are shown in first and last cases of (A.12). Partial degradation in the transient pulse is presented in the other equations. These equations were derived as follows.

$$\tau_n = \tau_g \left(k - \ln \left(1 - \frac{\tau_{n+1}}{\tau_g(k+1)} \right) \right) \quad (6.13)$$

Situations where τ_n is bigger than $k\tau_g$ and smaller than $(k+1)\tau_g$ are treated by (A.13).

$$\tau_n = \frac{\tau_{n+1} + \sqrt{\tau_{n+1}^2 + 4\tau_g^2}}{2} \quad (6.14)$$

Propagation cases where τ_n is bigger than $(k+1)\tau_g$ and smaller than $(k+3)\tau_g$ are treated by (A.14).

Function `sizeTransistors( s, g,  $\tau_n$  )` finds the smallest transistors width for a gate g according to the transient pulse duration τ_n . In Figure A.16, NMOS transistors of the first gate are “ON” and these transistors are the main responsible by the SET duration. Thus, only the NMOS transistors are sized in order to reduce the resistances $r1$ and $r2$ (the capacitance is indirectly increased as result of the diffusion areas).

The transistor sizing is modeled based on equations presented in Section A.4.2. The algorithm consists on applying the *bisection method* [Wei07a] to find the width of each pull-up and pull-down transistor of the gate.

6.6 Results

Table 6.3: The proposed transistor sizing to single event transient attenuation. Results show the area, timing and average power overhead for symmetric and asymmetric sizing techniques for particles with charge $Q = 0.3pC$ with final sensitivity of 50%.

Combinational Circuit	Sizing Methodology	Overhead		
		Area (%)	Power (%)	Timing (%)
C432	Symmetric	47.4	63.8	0.0
	Asymmetric	35.5	50.7	2.0
C880	Symmetric	88.0	72.4	0.0
	Asymmetric	69.2	51.6	0.0
C1355	Symmetric	62.4	38.6	16.0
	Asymmetric	50.6	29.5	15.8
C1908	Symmetric	47.0	35.5	12.0
	Asymmetric	37.0	29.0	8.8
Average overhead	Symmetric	61.2	52.7	7.0
	Asymmetric	48.0	40.2	6.65

Table 6.4: The proposed transistor sizing to single event transient attenuation. Results show the area, timing and average power overhead for symmetric and asymmetric sizing techniques for particles with charge $Q = 0.3pC$ with final sensitivity of 0%.

Combinational Circuit	Sizing Methodology	Overhead		
		Area (%)	Power (%)	Timing (%)
C432	Symmetric	69.8	105	1.2
	Asymmetric	50.0	59.7	0.0
C880	Symmetric	115.3	88.7	12.3
	Asymmetric	86.9	59.1	13.2
C1355	Symmetric	80.0	61.6	24.8
	Asymmetric	58.6	37.2	17.1
C1908	Symmetric	69.2	20.89	13.0
	Asymmetric	49.2	17.4	10.16
Average overhead	Symmetric	83.5	69.0	12.82
	Asymmetric	61.1	43.3	10.11

Table A.9 and A.10 show some results obtained by the proposed transistor sizing strategy. Results include a comparison between symmetric and asymmetric sizing methodologies for a $180nm$ technology process [ZC07]. The transient pulse propagation parameter

k was defined by hspice simulations as 0.8 for this technology. The transistor sizing was done aiming at reducing the sensitivity to 50% sensitivity (Table A.9) and 0% (Table A.10).

As discussed in Section 4.2.1.2, a study presented in [ZM06] shows that the deposited charge of very few particles is higher than $0.3pC$ for $180nm$ technologies at the atmosphere. We use this value in our experiments by considering as the worst case deposited charge.

The first important point shown by these results is the small overhead presented by the proposed methodology. The worst case was a 87% area overhead for complete protection (0% sensitivity) against particles with charge $Q = 0.3pC$. Results show an average 83% area overhead for the symmetric sizing and 61% for the asymmetric sizing. Power consumption presents 70% average overhead for the symmetric sizing against 43% for the asymmetric. Results show small timing penalties of 10% for the circuit with 0% sensitivity.

The asymmetric transistor sizing resulted in smaller area, power consumption and timing in comparison with the symmetric sizing. Despite of the penalties when designing radiation hardened circuits, results show the asymmetric sizing efficacy.

Table 6.5: A comparison among TMR, CWSP and the proposed sizing techniques. Results show area overhead and timing penalties to protect some circuits against particles with charge $Q = 0.3pC$ with final sensitivity of 0%.

Circuit	TMR		CWSP		Sizing	
	Area (%)	Timing (%)	Area (%)	Timing (%)	Area (%)	Timing (%)
C432	209.8	0.0	23.7	105.9	50.0	0.0
C880	213.8	0.0	44.6	83.4	86.9	13.2
C1355	222.3	0.0	36.6	106.1	58.6	17.1
C1908	219.2	0.0	36.5	64.9	49.2	10.1
Average	216.2	0.0	35.3	90.0	61.1	10.1

Table 6.5 shows a comparison among TMR, CWSP and the proposed sizing technique. It is possible to verify the high area overhead of the TMR technique. On the other hand, timing penalty is insignificant for the TMR due to the inclusion of only a voter in the critical path.

The CWSP technique presents small area overhead but timing penalties are very high due to the insertion of delay blocks. For particles with critical charge $Q_c = 0.3pC$, the transient duration is around $500ps$. A CWSP element with a delay bigger than $500ps$ is necessary to filter the transient.

For the presented combinational circuits, delay varies between $400ps$ and $700ps$. For this reason, the additional delay needed to attenuate the transient significantly increases the timing of the circuits. The area overhead of these combinational circuits is a function of the number of POs.

The proposed sizing technique presents better results in comparison with the TMR and CWSP techniques. The asymmetric sizing technique guarantees smaller area overhead

due to selective transistor sizing as well as smaller timing penalties because new elements are not inserted in the critical paths.

6.7 Conclusions

A new transistor sizing algorithm aiming at protecting combinational logic circuit against single event transients is presented in this Chapter. The sensitivity of the circuit is analyzed by taking into account the logical and electrical masking.

The proposed technique consists on sizing only transistors directly related to the SET attenuation. It is known that transistors PMOS and NMOS have different characteristics in relation to mobility, impurities concentration and, as consequence, the delay. For a given particle with charge Q , PMOS and NMOS transistors present different attenuation characteristics. Thus, the model considers independently pull-up and pull-down blocks.

Besides, the model takes into account propagation characteristics in which the degradation of the transient pulse is considered in order to reduced sizing penalties.

The importance of this method is presented in the results. Results show smaller area, timing and power consumption overhead in comparison with a symmetrical methodology. The reduced timing penalties presented by the sizing methodology allows the development of high frequency circuits, with low overhead concerning area and power.

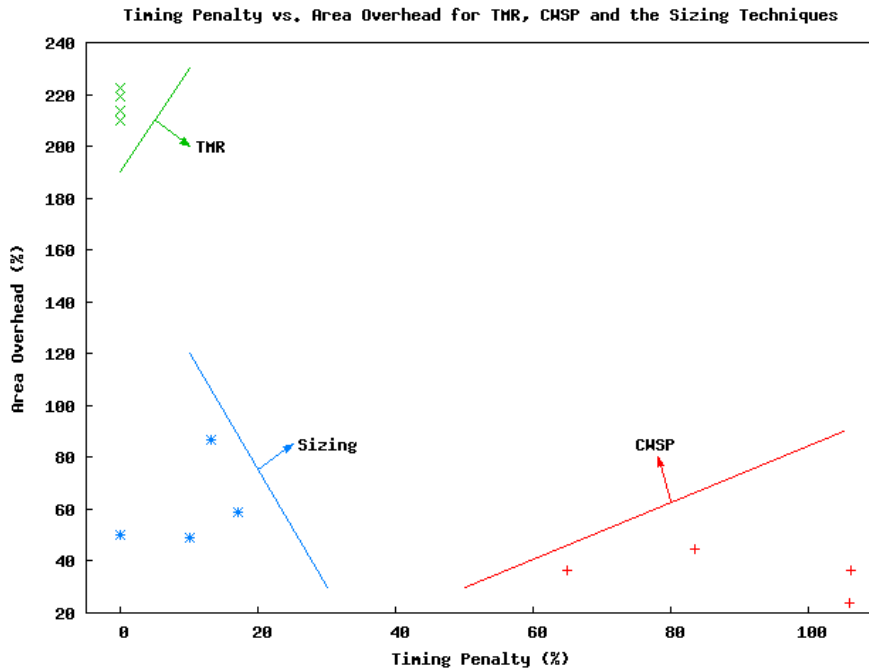


Figure 6.6: Timing penalty versus area overhead for TMR, CWSP and the proposed sizing technique.

Figure 6.6 summarizes TMR, CWSP and the proposed sizing techniques. Penalties of the proposed transistor sizing technique are grouped in the bottom-left corner. This characteristic highlights the efficiency of the proposed technique. The TMR technique fails because of the huge area overhead, while the CWSP technique presents big timing

penalties.

7 CONCLUSION

The contributions of this thesis are basically divided in two major parts. The first is related to the development of a new methodology able to generate optimized circuits concerning timing and power consumption. The transistor level design flow, as it is called, optimizes every single gate of a circuit according to the capacitances in which it is involved.

The advent of the deep submicron technologies has included a myriad of new challenges in the design of circuits. The geometries are shrinking, power supplies are getting lower and the logic density is reaching a very high rate. In addition, the increased number of metal layers, associated with these deep submicron characteristics, are shifting the design paradigm from circuits with the delay dominated by the logic to circuits with the delay dominated by interconnections.

Process variability and the huge variance in the capacitance per unit length emphasize the need of a transistor level optimization because it is practically impossible to predict interconnects before the layout phase.

The proposed transistor level design flow is based on academic and commercial tools. Commercial tools include logic synthesis, placement & routing, and switch level static timing & power consumption analysis. Academic tools were developed to cope with the gaps between the conventional standard cell and transistor level methodologies.

Basically, the transistor level design flow presents four differences in comparison with the standard cell methodology:

1. *Library generation:* The layout of cells is not generated at library generation time. This allows significantly to increase the number of logic functions.
2. *Transistor level optimization:* Transistor level optimization allows to find optimized transistors width according to the capacitances involved with the gate.
3. *Layout generation:* The layout generation is performed after the transistor optimization to cope with a very large range of possibilities concerning transistors width.
4. *A feedback flow:* A flow with feedback allows to take into account the extracted capacitances during the transistor optimization process.

A power leakage reduction method is also proposed in which a gate length biasing technique is applied in order to reduce the static current. The gate length biasing tech-

nique consists on adjusting the length of transistors aiming at reducing the current flowing through it.

All these features allow to mitigate the timing closure problem. Results show that this methodology is very promising. Comparisons between the transistor level methodology and the standard cell approach show interesting results where the proposed methodology presents around 11% of delay improvement and more than 30% power saving.

The second contribution of this thesis concerns the application of the transistor level design flow in the protection of integrated circuits against single event effects (SEE). The main aspect concerning the transistor level design flow is the possibility to develop a new transistor sizing methodology targeting radiation-hardened combinational circuits.

Technology scaling also has effects in integrated circuits concerning functional failure due to SEEs. Gate length reduction and low supply voltages make circuits sensitive to energetic particles that were worthless in older technologies.

Two contributions are presented in respect to the protection of circuit against SEE. The first is related with the insertion of timing redundancy in sequential elements in order to cope with single event. The main idea consists on applying the concepts proposed by Anghel in [Ang00] about the Code Word State Preserving (CWSP) technique in latches and flipflops. The technique was applied to some microprocessors and results show a worst case of 85% area overhead and 7% timing penalty.

The second contribution related with single event effects protection involves a new transistor sizing methodology. The sensitivity of combinational circuits are defined by logical and electrical masking analysis. The sensitivity of each node allows to size each individual gate as a function of the particle charge.

The logical masking gives the probability of a transient pulse to be masked by the circuit logic, which is computed by controllability and observability techniques. The electrical masking describes whether a transient pulse in a node is not propagated to the primary outputs due to electrical attenuation.

The sizing method is based on an analytical model of transient pulse detection presented in [WVK07, WVNK07]. The model is based on electrical device parameters. A propagation method is used, where the gate delay and the transient pulse duration are evaluated. The propagation model is very important because it assumes that the transient pulse must be attenuated during the whole path and not in the struck node. Thus, oversized transistors are avoided. The model also considers independently pull-up and pull-down blocks. Only transistors directly related to the SET attenuation are sized.

The proposed sizing technique presented small area, timing and power consumption overhead in comparison with the TMR and CWSP techniques, allowing the development of high frequency circuits, with low area and power overhead. Results show an average 61% area overhead, 43% power consumption increase and 10% delay penalty.

REFERENCES

- [AC99] Sekan Askar and Maciej Ciesielski. Analytical approach to custom datapath design. In *IEEE/ACM International Conference on Computer-Aided Design, Digest of Technical Papers*, pages 98–101, 1999.
- [adLK03] Fernanda Gusm ao de Lima Kastensmidt. *Designing Single Event Upset Mitigation Techniques for Large SRAM-Based FPGA Components*. PhD thesis, Instituto de Informatica, UFRGS, Porto Alegre, 2003.
- [AHH⁺95] C. J. Alpert, T. C. Hu, J. H. Huang, A. B. Kahng, and D. Karger. Prim-dijkstra tradeoffs for improved performance-driven routing tree design. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 14:890–896, 1995.
- [AND98] C. ANDREW. Vlsi datapath choices: Cell-based versus fullcustom, 1998.
- [Ang00] Lorena Anghel. *Les Limites Technologiques du Silicium et Tolerance aux Fautes*. PhD thesis, Institut Polytechnique de Grenoble, France, 2000.
- [Bal00] Florin Balasa. Modeling non-slicing floorplans with binary trees. In *IC-CAD*, pages 13–16, 2000.
- [Bal01] Florin Balasa. Device-level placement for analog layout: an opportunity for non-slicing topological representations. In *ASP-DAC '01: Proceedings of the 2001 conference on Asia South Pacific design automation*, pages 281–286, New York, NY, USA, 2001. ACM Press.
- [Bas06] Rodrigo Possamai Bastos. Design a robust microprocessor to soft errors. Master's thesis, Instituto de Informatica, UFRGS, Porto Alegre, 2006.
- [Bau05] Robert Baumann. Soft errors in advanced computer systems. *IEEE Des. Test*, 22(3):258–266, 2005.
- [BBMM04] Pietro Babighian, Luca Benini, Alberto Macii, and Enrico Macii. Post-layout leakage power minimization based on distributed sleep transistor insertion. In *ISLPED '04: Proceedings of the 2004 international symposium on Low power electronics and design*, pages 138–143, New York, NY, USA, 2004. ACM Press.

- [BBR02] Anil Bahuman, Benjamin Bishop, and Khaled Rasheed. Automated synthesis of standard cells using genetic algorithms. In *Proceedings on the IEEE Computer Society Annual Symposium on VLSI*, pages 126–133, 2002.
- [BC05] L. Behjat and A. Chiang. Fast integer linear programming based models for vlsi global routing. In *Proceedings of the IEEE International Symposium on Circuits and Systems, 2005. ISCAS 2005.*, pages 6238–6243, 2005.
- [BCV06] Sarvesh Bhardwaj, Yu Cao, and Sarma Vrudhula. Statistical leakage minimization through joint selection of gate sizes, gate lengths and threshold voltage. In *Proceedings of the 2006 conference on Asia South Pacific design automation. ASP-DAC '06*, pages 953–958, New York, NY, USA, 2006. ACM Press.
- [Ber06] Michel Berkelaar. Ip solve 5.5.0.9. Technical report, 2006.
- [BLGR04] F. Bastian, C. Lazzari, J. Guntzel, and R. Reis. A new transistor folding algorithm applied to an automatic full-custom layout generation tool. In *Proceedings of the 14th International Workshop on Power and Timing Modeling*, pages 732–741, 2004.
- [BMK02] Florin Balasa, Sarat C. Maruvada, and Karthik Krishnamoorthy. Efficient solution space exploration based on segment trees in analog placement with symmetry constraints. In *Proceedings of the 2002 IEEE/ACM international conference on Computer-aided design*, pages 497–502, New York, NY, USA, 2002. ACM Press.
- [BMK04] Florin Balasa, Sarat C. Maruvada, and Karthik Krishnamoorthy. On the exploration of the solution space in analog placement with symmetry constraints. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 23(2):177–191, February 2004.
- [BV94] D. Bessot and R. Velazco. Design of seu-hardened cmos memory cells: The hit cell. In *Proceedings of the 1994 RADECS Conference*, pages 563–570, 1994.
- [BVR06] Laleh Behjat, Anthony Vannelli, and William Rosehart. Integer linear programming models for global routing. *INFORMS J. on Computing*, 18(2):137–150, 2006.
- [CAD07a] CADENCE. Global synthesis for timing closure - the impact on design closure. Technical report, 2007.
- [CAD07b] CADENCE. <http://www.cadence.com>. April 2007.
- [CAL02] Maciej Ciesielski, Sekan Askar, and Samuel Levitin. Analytical approach to layout generation of datapath cells. In *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, number 12, pages 1480–1488, Dec 2002.

- [CCS05] Tony Chan, Jason Cong, and Kenton Sze. Multilevel generalized force-directed method for circuit placement. In *ISPD '05: Proceedings of the 2005 international symposium on Physical design*, pages 185–192, New York, NY, USA, 2005. ACM Press.
- [Che06] C.K. Cheng. Timing closure using layout based design process. 2006.
- [CHP00] Wei CHEN, Cheng-Ta HSIEH, and Massoud PEDRAM. Simultaneous gate sizing and fanout optimization. In *Proceedings of the IEEE/ACM international conference on Computer-aided design*, pages 374–378, San Jose, California, 2000.
- [CK02] David Chinnery and Kurt Keutzer. *Closing the Gap Between ASIC & Full Custom: Tools and Techniques for High-Performance ASIC Design*. Kluwer Academic Publisher Group, 2002.
- [CNV96] T. Calin, M. Nicolaidis, and R. Velazco. Upset hardened memory design for submicron cmos technology. *IEEE Transactions on Nuclear Science*, 43:2874–2878, December 1996.
- [CRO⁺05] J. M. Cazeaux, D. Rossi, M. Omana, C. Metra, and A. Chatterjee. On transistor level gate sizing for increased robustness to transient faults. In *IOLTS '05: Proceedings of the 11th IEEE International On-Line Testing Symposium*, pages 23–28, Washington, DC, USA, 2005. IEEE Computer Society.
- [CS00] Jason Cong and Majid Sarrafzadeh. Incremental physical design. In *Proceedings of the International Symposium on Physical design*, pages 84–92, San Diego, California, USA, 2000. New York: ACM Press.
- [DDCS04] Yuvraj S. Dhillon, Abdulkadir U. Diril, Abhijit Chatterjee, and Adit D. Singh. Sizing cmos circuits for increased transient error tolerance. In *IOLTS '04: Proceedings of the International On-Line Testing Symposium, 10th IEEE (IOLTS'04)*, page 11, Washington, DC, USA, 2004. IEEE Computer Society.
- [DGC⁺92] C. Dufour, P. Garnier, T. Carriere, J. Beaucour, R. Ecoffet, and M. Labrunee. Heavy ion induced single hard errors on submicronic memories. *IEEE Transactions on Nuclear Science*, pages 1693 – 1697, 1992.
- [DKCP94] A. Dharchoudhury, S. M. Kang, H. Cha, and J. H. Patel. Fast timing simulation of transient faults in digital circuits. In *ICCAD '94: Proceedings of the 1994 IEEE/ACM international conference on Computer-aided design*, pages 719–722, Los Alamitos, CA, USA, 1994. IEEE Computer Society Press.
- [DM03] P.E. Dodd and L.W. Messengill. Basic mechanisms and modeling of single-event upset in digital microelectronics. In *Proceedings of the IEEE Transactions on Nuclear Science*, volume 50, pages 583–602, June 2003.

- [DRSVW87] E. Detjens, R. Rudell, A.L. Sangiovanni-Vincentelli, and A. Wang. Technology mapping in mis. In *Proceedings on the ICCAD*, pages 116–119, 1987.
- [ELEHS03] Henrik Eriksson, Per Larsson-Edefors, Tomas Henriksson, and Christer Svensson. Full-custom vs. standard-cell design flow: an adder case study. In *ASPDAC: Proceedings of the 2003 conference on Asia South Pacific design automation*, pages 507–510, New York, NY, USA, 2003. ACM Press.
- [GCY99] Pei-Ning Guo, Chung-Kuan Cheng, and Takeshi Yoshimura. An o-tree representation of non-slicing floorplan and its applications. In *DAC*, pages 268–273, 1999.
- [GH00] A. GUPTA and J. P. HAYES. Clip: integer-programming-based optimal layout synthesis of 2d cmos cells. In *Proceedings of the IEEE Transactions on Design Automation of Electronic Systems*, volume 5, pages 510–547, New York, NY, USA, 2000. New York: ACM Press.
- [GKSS04] Puneet Gupta, Andrew B. Kahng, Puneet Sharma, and Dennis Sylvester. Selective gate-length biasing for cost-effective runtime leakage control. In *Proceedings of the 41st annual conference on Design automation. DAC '04*, pages 327–330, New York, NY, USA, 2004. ACM Press.
- [GMD⁺97] Mohan Guruswamy, Robert L. Maziasz, Daniel Dulitz, Srilata Raman, Venkat Chiluvuri, Andrea Fernandez, and Larry G. Jones. CELLERITY: A fully automatic layout synthesis system for standard cell libraries. In *Proceedings of the 34th Design Automation Conference*, pages 327–332, 1997.
- [GR01] Prakash Gopalakrishnan and Rob A. Rutenbar. Direct transistor-level layout for digital blocks. In *ICCAD '01: Proceedings of the 2001 IEEE/ACM international conference on Computer-aided design*, pages 577–584, Piscataway, NJ, USA, 2001. IEEE Press.
- [GSB⁺04] M.J. Gadlage, R.D. Schrimpf, J.M. Benedetto, P.H. Eaton, D.G. Mavis, M.Sibley, K. Avery, and T.L. Turflinger. Single event transient pulse widths in digital microcircuits. *IEEE Transactions on Nuclear Science*, 51:3285 – 3290, December 2004.
- [HCR⁺03] Sung-Woo Hur, Tung Cao, Karthik Rajagopal, Yegna Parasuram, Amit Chowdhary, Vladimir Tiourin, and Bill Halpin. Force directed mongrel with physical net constraints. In *DAC '03: Proceedings of the 40th conference on Design automation*, pages 214–219, New York, NY, USA, 2003. ACM Press.
- [Hen02] Renato Hentschke. Algoritmos para o posicionamento de celulas em circuitos vlsi. Master's thesis, Instituto de Informatica, UFRGS, Porto Alegre, 2002.

- [Hen07] Renato Hentschke. *Algorithms for Wire Length Improvement of VLSI Circuits With Concern to Critical Paths*. PhD thesis, Instituto de Informatica, UFRGS, Porto Alegre, 2007.
- [HFPR06] Renato Hentschke, Guilherme Flach, Felipe Pinto, and Ricardo Reis. Quadratic placement for 3d circuits using z-cell shifting, 3d iterative refinement and simulated annealing. In *SBCCI '06: Proceedings of the 19th annual symposium on Integrated circuits and systems design*, pages 220–225, New York, NY, USA, 2006. ACM Press.
- [HN06] Tino Heijmen and Andre Nieuwland. Soft-error rate testing of deep-submicron integrated circuits. In *ETS '06: Proceedings of the Eleventh IEEE European Test Symposium (ETS'06)*, pages 247–252, Washington, DC, USA, 2006. IEEE Computer Society.
- [HNJR07] Renato F. Hentschke, Jaganathan Narasimham, Marcelo O. Johann, and Ricardo L. Reis. Maze routing steiner trees with effective critical sink optimization. In *Proceedings of the 2007 international symposium on Physical design. ISPD '07*, pages 135–142, New York, NY, USA, 2007. ACM Press.
- [HO01] Masanori Hashimoto and Hidetoshi Onodera. Post-layout transistor sizing for power reduction in cell-based design. In *ASP-DAC '01: Proceedings of the 2001 conference on Asia South Pacific design automation*, pages 359–365, New York, NY, USA, 2001. ACM Press.
- [HSDU⁺90] R. Harboe-Sorensen, E.J. Daly, C.I. Underwood, J. Ward, and L. Adams. The behaviour of measured seu at low altitude during periods of high solar activity. *IEEE Transactions on Nuclear Science*, pages 1938 – 1946, 1990.
- [IEE02] IEEE Design & Test Staff. Deep-submicron challenges. *IEEE Des. Test*, 19(2):3, 2002.
- [JJ05] Samir Jasarevic and Goran Jerin. Automated soft error analysis and protection in combinational logic circuits. Master's thesis, Lund University, Sweden, 2005.
- [Jun07] Adriel Mota Ziesemer Junior. Geracao automatica de partes operativas de circuitos vlsi. Master's thesis, Instituto de Informatica, UFRGS, Porto Alegre, 2007.
- [KAB⁺03] Nam Sung Kim, Todd Austin, David Blaauw, Trevor Mudge, Krisztián Flautner, Jie S. Hu, Mary Jane Irwin, Mahmut Kandemir, and Vijaykrishnan Narayanan. Leakage current: Moore's law meets static power. *Computer*, 36(12):68–75, 2003.
- [KMS05] Andrew B. Kahng, Swamy Muddu, and Puneet Sharma. Defocus-aware leakage estimation and control. In *Proceedings of the 2005 international symposium on Low power electronics and design. ISLPED '05*, pages 263–268, New York, NY, USA, 2005. ACM Press.

- [KSJA91] Jurgen Kleinhans, Georg Sigl, Frank Johannes, and Kurt Antreich. Gordian: Vlsi placement by quadratic programming and slicing optimization. In *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, volume 10, pages 356–365, March 1991.
- [LAR05a] C. LAZZARI, L. ANGHEL, and R. REIS. On implementing a soft error hardening technique by using an automatic layout generator: Case study. In *11th IEEE International On-Line Testing Symposium*, pages 29–34, 2005.
- [LAR05b] C. Lazzari, L. Anghel, and R. Reis. A transistor placement technique using genetic algorithm and analytical programming. In *Proceedings of the IFIP WG.5 Conference on Very Large Scale Integration System-on-Chip*, pages 559–564, 2005.
- [Laz03] Cristiano Lazzari. Automatic layout generation of static cmos circuits targeting delay and power reduction. Master’s thesis, Instituto de Informatica, UFRGS, Porto Alegre, 2003.
- [LBM⁺00] K. A. LaBel, C.E. Barnes, P.W. Marshall, C.J. Marshall, A.H. Johnston, R.A. Reed, J.L. Barth, C.M. Seidleck, S.A. Kayali, and M.V. O’Byran. A roadmap for nasa’s radiation effects research in emerging microelectronics and photonics. In *proceedings of the 2000 IEEE Aerospace Conference*, volume 5, pages 535–545, 2000.
- [LCLH96] John Lillis, Chung-Kuan Cheng, Ting-Ting Y. Lin, and Ching-Yen Ho. New performance driven routing techniques with explicit area/delay trade-off and simultaneous wire sizing. In *Proceedings of the 33rd annual conference on Design automation. DAC ’96*, pages 395–400, New York, NY, USA, 1996. ACM Press.
- [LDGR03] Cristiano Lazzari, Cristiano V. Domingues, Jose Luis Guntzel, and Ricardo A. L. Reis. A new macro-cell generation strategy for three metal layer cmos technologies. In *Proceedings of the VLSI-Soc*, pages 143–147, Darmstadt, Germany, 2003.
- [LH03] Changbo Long and Lei He. Distributed sleep transistor network for power reduction. In *Proceedings of the 40th conference on Design automation. DAC ’03*, pages 181–186, New York, NY, USA, 2003. ACM Press.
- [LH05] Yan Lin and Lei He. Leakage efficient chip-level dual-vdd assignment with time slack allocation for fpga power reduction. In *Proceedings of the 42nd annual conference on Design automation. DAC ’05*, pages 720–725, New York, NY, USA, 2005. ACM Press.
- [MAB⁺03] Howard Matis, Gordon Aubrecht, A. Baha Balantekin, Wolfgang Bauer, John Beacom, Elizabeth J. Beise, David Bodansky, Edgardo Browne, Peggy Carlock, and Yuen-Dat Chan. *Space Applications: Radiation-Induced Effects*. Available at <http://www.lbl.gov/abc/wallchart/guide.html>, 2003.

- [ME95] Larry McMurchie and Carl Ebeling. Pathfinder: a negotiation-based performance-driven router for fpgas. In *Proceedings of the 1995 ACM third international symposium on Field-programmable gate arrays. FPGA '95*, pages 111–117, New York, NY, USA, 1995. ACM Press.
- [ME02] D.G. Mavis and P.H. Eaton. Soft error rate mitigation techniques for modern microcircuits. In *Proceedings of the 40th International Reliability Physics Symposium*, pages 216–255, 2002.
- [Mes82] G Messenger. Collection of charge on junction nodes from ion tracks. In *Proceedings of the IEEE Transactions on Nuclear Science*, pages 2024–2031, 1982.
- [MG88] Sivanarayana Mallela and Lov K. Grover. Clustering based simulated annealing for standard cell placement. In *DAC '88: Proceedings of the 25th ACM/IEEE conference on Design automation*, pages 312–317, Los Alamitos, CA, USA, 1988. IEEE Computer Society Press.
- [Moo03] Gordon E. Moore. No exponential is forever: but “forever” can be delayed! In *Digest of Technical Papers. ISSCC. 2003 IEEE International Solid-State Circuits Conference*, volume 1, pages 20–23, 2003.
- [MT03] Kartik Mohanram and Nur A. Touba. Cost-effective approach for reducing soft error failure rate in logic circuits. *etc*, 00:893, 2003.
- [MTB00] Fan Mo, Abdallah Tabbara, and Robert Brayton. A force-directed macro-cell placer. In *IEEE/ACM International Conference on Computer Aided Design, ICCAD-2000*, pages 177–180, 2000.
- [NC97] M. Nicolaidis and T. Calin. A theory of perturbation tolerant asynchronous fsm and its application on the design of perturbation tolerant memories. *1997 European Test Workshop*, pages –, May 1997.
- [NDB⁺02] Siva Narendra, Vivek De, Shekhar Borkar, Dimitri Antoniadis, and Anantha Chandrakasan. Full-chip sub-threshold leakage power prediction model for sub-0.18 μ m cmos. In *ISLPED '02: Proceedings of the 2002 international symposium on Low power electronics and design*, pages 19–23, New York, NY, USA, 2002. ACM Press.
- [NJJ06] Andre K. Nieuwland, Samir Jasarevic, and Goran Jerin. Combinational logic soft error analysis and protection. In *IOLTS '06: Proceedings of the 12th IEEE International Symposium on On-Line Testing*, pages 99–104, Washington, DC, USA, 2006. IEEE Computer Society.
- [Nor01] Eugene Normand. Single event upset at ground level. *IEEE Transactions on Nuclear Science*, 43:2742 – 2750, December 2001.
- [NW07] James Noyes and Weisstein. Linear programming. From MathWorld – A Wolfram Web Resource. Available in <http://mathworld.wolfram.com/LinearProgramming.html>, 2007.

- [OCG02] Ralph Otten, Raul Camposano, and Patrick Groeneveld. Design automation for deepsubmicron - present and future. In *Proceedings of the Design, Automation and Test in Europe Conference and Exhibition*, pages 650–657, Washington, DC, USA, 2002. [S.l.]: IEE Computer Society.
- [PCBO94] C. Poivey, T. Carriere, J. Beaucour, and T.R. Oldham. Characterization of single hard errors (she) in 1 m-bit srams from single ion. *IEEE Transactions on Nuclear Science*, pages 2235 – 2239, 1994.
- [RBB05] Rob Roy, Debashis Bhattacharya, and Vamsi Boppana. Transistor-level optimization of digital designs with flex cells. *IEEE Computer Society*, 38:53–61, Feb 2005.
- [RC06] Sherief Reda and Amit Chowdhary. Effective linear programming based placement methods. In *ISPD '06: Proceedings of the 2006 international symposium on Physical design*, pages 186–191, New York, NY, USA, 2006. ACM Press.
- [Roc88] L. Rockett. An seu hardened cmos data latch design. *IEEE Transaction on Nuclear Science*, NS-35(6):1682–1687, Dec 1988.
- [RRAR97] Andre I. Reis, Ricardo A. L. Reis, D. Auverne, and M. Robert. Library free technology mapping. *VLSI: Integrated Systems on Silicon, IFIP TC10 WG10.5 International Conference in Very Large Scale Integration*, pages 303–314, August 1997.
- [RS99] Michael A. Riepe and Karem A. Sakallah. Transistor level micro-placement and routing for two-dimensional digital vlsi cell synthesis. In *ISPD '99: Proceedings of the 1999 international symposium on Physical design*, pages 74–81, New York, NY, USA, 1999. ACM Press.
- [RS03] Michael A. Riepe and Karem A. Sakallah. Transistor placement for non-complementary digital vlsi cell synthesis. *ACM Trans. Des. Autom. Electron. Syst.*, 8(1):81–107, 2003.
- [RTL95] Sanjay Rekhi, J. Donald Trotter, and Daniel H. Linder. Automatic layout synthesis of leaf cells. In *DAC '95: Proceedings of the 32nd ACM/IEEE conference on Design automation*, pages 267–272, New York, NY, USA, 1995. ACM Press.
- [Sec88] Carl Sechen. Chip-planning, placement, and global routing of macro/custom cell integrated circuits using simulated annealing. In *25th ACM/IEEE Proceedings of the Design Automation Conferenc*, pages 73–80, 1988.
- [SIA97] SIA. The national technology roadmap for semiconductors. Semiconductor Industry Association, 1997.

- [SS01] T. Serdar and C. Sechen. Automatic datapath tile placement and routing. In *Proceedings of the Design, Automation and Test in Europe*, pages 13–16, 2001.
- [SWL⁺03] Cristiano Santos, Gustavo Wilke, Cristiano Lazzari, Jose Luis Guntzel, and Ricardo Reis. A transistor sizing method applied to an automatic layout generation tool. In *Proceedings of the SYMPOSIUM ON INTEGRATED CIRCUITS AND SYSTEMS DESIGN*, pages 303–307, 2003.
- [SYN07] SYNOPSYS. <http://www.synopsys.com>. April 2007.
- [Tay03] Satoshi Tayu. A simulated annealing approach with sequence-pair encoding using a penalty function for the placement problem with boundary constraints. In *ASPDAC: Proceedings of the 2003 conference on Asia South Pacific design automation*, pages 319–324, New York, NY, USA, 2003. ACM Press.
- [VC04] Natarajan Viswanathan and Chris Chong-Nuen Chu. Fastplace: efficient analytical placement using cell shifting, iterative local refinement and a hybrid net model. In *ISPD '04: Proceedings of the 2004 international symposium on Physical design*, pages 26–33, New York, NY, USA, 2004. ACM Press.
- [VWSS04] Miodrag Vujkovic, David Wadkins, Bill Swartz, and Carl Sechen. Efficient timing closure without timing driven placement and routing. In *DAC '04: Proceedings of the 41st annual conference on Design automation*, pages 268–273, New York, NY, USA, 2004. ACM Press.
- [WCL91] S. Whitaker, J. Canaris, and K. Liu. Seu hardened memory cells for a cclds reed solomon encoder. *IEEE Transaction on Nuclear Science*, NS-36(6):1471–1477, December 1991.
- [WE93] Neil Weste and Kamran Esraghian. *Principles of CMOS VLSI Design - A Systems Perspective*. Addison-Wesley Publishing Company, 1993.
- [Wei07a] Eric W. Weisstein. Bisection. From MathWorld – A Wolfram Web Resource. Available in <http://mathworld.wolfram.com/Bisection.html>, 2007.
- [Wei07b] Eric W. Weisstein. Euler path. From MathWorld – A Wolfram Web Resource. Available in <http://mathworld.wolfram.com/EulerPath.html>, 2007.
- [WLS05] Qingzhou Wang, John Lillis, and Shubhankar Sanyal. An lp-based methodology for improved timing-driven placement, 2005.
- [WRJS74] David C. Wilson and II Robert J. Smith. An experimental comparison of force directed placement techniques. In *DAC '74: Proceedings of the 11th workshop on Design automation*, pages 194–199, Piscataway, NJ, USA, 1974. IEEE Press.

- [WVK07] G.I. Wirth, M.G. Vieira, and F.G. Lima Kastensmidt. Accurate and computer efficient modelling of single event transients in cmos circuits. *IET Circuits, Devices & Systems*, 1:137–142, April 2007.
- [WVNK07] G.I. Wirth, M.G. Vieira, Egas H. Neto, and F.G. Lima Kastensmidt. Modelling the sensivity of cmos circuits to radiation induced single event transients. *Microelectronics Reliability*, 47(3), March 2007.
- [ZC07] Wei Zhao and Yu Cao. Predictive technology model for nano-cmos design exploration. *J. Emerg. Technol. Comput. Syst.*, 3(1):1, 2007.
- [ZM06] Quming Zhou and K. Mohanram. Gate sizing to radiation harden combinational logic. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 25:155– 166, Jan 2006.

APPENDIX A GENERATION AUTOMATIQUE DE CIRCUITS DURCIS AUX RAYONNEMENTS AU NIVEAU TRANSISTOR

A.1 Introduction

Les technologies submicroniques profondes (DSM, de l'anglais *Deep Submicron*) ont inséré des nouveaux défis dans le projet de circuits intégrés à cause de la réduction des géométries, la réduction de la tension d'alimentation, l'augmentation de la fréquence et la densité élevée de la logique [CS00].

Les interconnexions présente une importance très élevée en technologies DSM. Les nouvelles technologies ont décalé le paradigme de projet d'un dominé pour la logique pour un projet dominé par l'interconnexion [Che06].

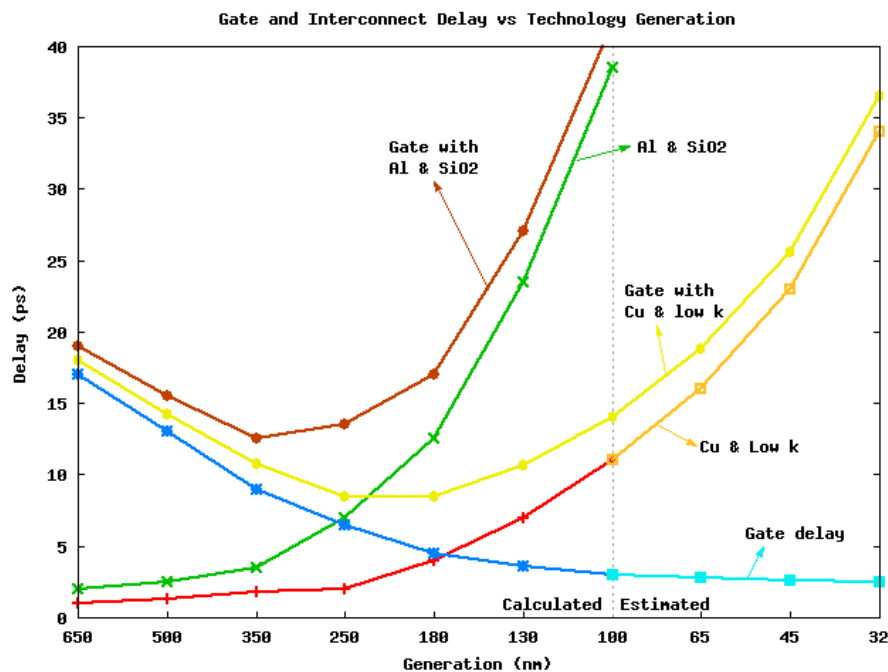


Figure A.1: retard des portes et des interconnexion versus les génération de technologie [SIA97]. Les retards pour les technologies au dessous de 100nm sont estimés.

La figure A.1 présente le retard comme fonction des générations de technologie [SIA97]. Technologies au dessous de $100nm$ sont estimées. Il est possible de remarquer que le retard des interconnexions excède les retard des portes dans la technologie $250nm$ si Al est utilisé pour les interconnexions et SiO_2 est utilisé comme diélectrique.

Le retard des interconnexions présente plus d'importance dans le processus de technologie $180nm$ quand le Cu est employé dans les interconnexions et un diélectrique *low k* est utilisé comme isolateur.

Le retard des interconnexions n'est pas le seul facteur important pour les technologies DSM. La densité élevée de la logique et le plus grand nombre de couche en métal font de la prévision des interconnexions un très compliqué tâche. La capacité par unité de longueur peut varié autour 35 fois en technologies $180nm$, par exemple [VWSS04].

La logique est optimisée pour le retard, la surface occupé et la puissance avec des capacités assumées, mais leurs valeurs réelles ne sont pas connues jusque la phase de génération du layout. *Timing closure* ne peut pas être réalisée avec cette prévision imprécise des interconnexions et sans l'intégration du layout avec le processus entier.

L'arrivée des technologies DSM fait également obligatoire le contrôle de la puissance statique. Les concepteurs de circuits intégrés ont concerné la puissance dynamique pendant toute l'évolution technologique. Autrement, la fuite de puissance n'était pas prise en considération en technologies plus anciennes à cause de sa basse importance.

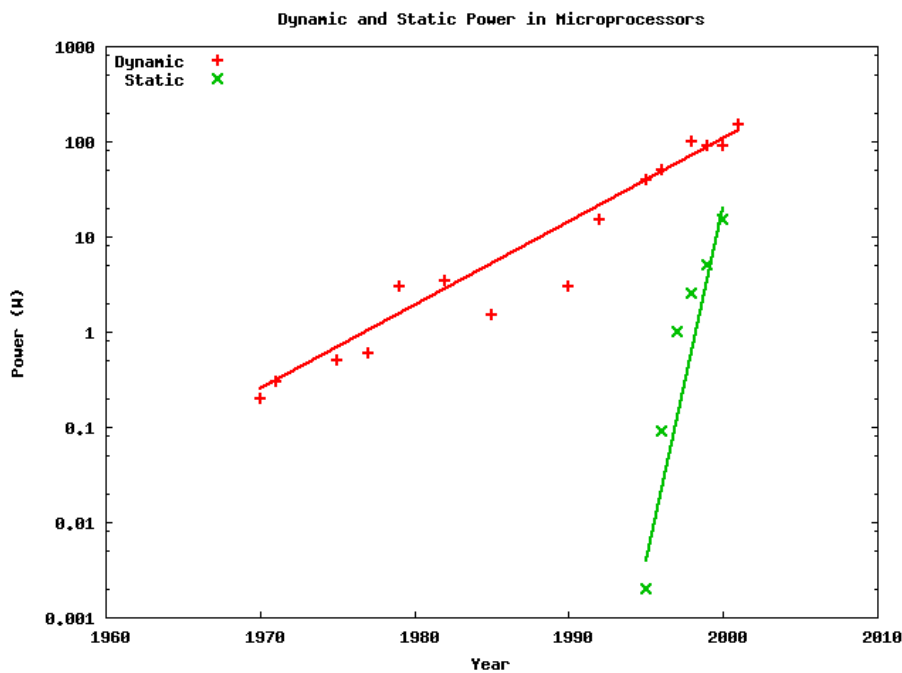


Figure A.2: Active and static power in microprocessors [Moo03].

La figure A.2 présente la puissance dynamique et la puissance statique dans le microprocesseur comme fonction de l'année [Moo03]. On projette que la fuite de puissance excède la puissance dynamique pour le technologies $65nm$ [KAB⁺03].

Ces défis submicroniques soulignent le besoin d'outils électroniques pour la concep-

tion automatisée (EDA, de l'anglais *Electronic Design Automation*) capables de générer et valider des circuits intégrés.

La conception basée dans les cellules standard ont dominé la génération de layout des circuits numériques VLSI dû à quelques vertus [GR01]. Les cellules standard cachent les plus désagréables détails sur les règles de dessin, les entrées et sorties sont accessibles facilement, les cellules sont facilement assemblées dans le circuit et les cellules sont caractérisées pour le retard et la puissance.

Cependant, quelques travaux ont proposé des modifications dans le flux traditionnel de projet afin de faire face à ces nouveaux défis submicroniques.

Un méthode d'optimisation après le layout propose une méthodologie de réduction de taille de transistor avec la préservation des interconnexions [HO01]. Le méthode a obtenu une réduction de puissance de 65% sans augmenter le délai du circuit.

Vujkovic *et al* a montré qu'une bibliothèque de cellules standard habituellement ne peut pas traiter le énorme variance dans les capacité avec un petit nombre de versions d'une cellule [VWSS04]. Autrement, si une cellule é capable de traiter la capacité exigée, la cellule dépense beaucoup de puissance dû aux transistors dimensionnés.

Le concept des cellules *flex* est présenté dans le papier [RBB05]. Ils proposent l'identification et l'optimisation d'un nombre minimum de cellules critiques visant augmentant la performance des circuits intégrés.

La première contribution de cette thèse est la proposition d'un nouveau paradigme dans la conception des circuits. Dans le flux traditionnel de conception, le layout des cellules sont produites et caractérisées pour le délai et pour la puissance. Ces cellules sont groupées dans une bibliothèque et utilisé pendant le projet du circuit intégré.

Nous proposons un flux de conception au niveau de transistor dans lequel la génération du layout est intégrée dans le processus de conception. Le flux de conception au niveau transistor consiste dans l'optimisation individuel de chaque porte d'un circuit. Ainsi, les transistors dans les chemins critiques du circuit sont optimisés, mais les portes dans les chemins non critiques sont maintenues avec transistors des tailles minimum.

La réduction de la puissance dynamique est une conséquence directe de l'optimisation des taille des transistors. L'optimisation au niveau transistor inclut une méthodologie d'optimisation de la longueur des portes afin de faire face à la fuite de courant. La longueur des transistors dehors les chemins critiques sont augmenté pour réduire la puissance statique sans pénalise le retard du circuit.

Les effets de rayonnement ont été également augmentés dans les technologies DSM. Les niveaux de dimension réduites des transistors et la basse de tension causent une augmentation de la taux d'erreur dans des circuits intégrés.

Dans le siècle 20, on a assumé que les SEEs (*Single Event Effects* dans la littérature anglaise) sont concernés principalement dans le espace. Ainsi, les techniques de durcissement on été appliqué à ces applications pour éviter la perte d'information ou erreur fonctionnel.

Cependant, le *scaling* des technologies réduit la fiabilité des circuits intégrés comme résultat des SEEs. L'énergie des neutrons, par exemple, varie entre 1mev à 10mev à l'atmosphère [Nor01]. Sans importance en technologies plus anciennes, l'énergie dans le flux atmosphérique de neutron est assez pour affecter rigoureusement la fonctionnalité des circuits intégrés courants.

Kastensmidt signale que les éléments de mémoire en technologie de $0.25\mu m$ et logique combinatoire composées de transistors avec la longueur plus petite que $0.13\mu m$ peuvent être sujets à SEE en circuits dans l'atmosphère [adLK03].

Pour cette raison, le développement d'outils électroniques de conception automatisée capables faire face au SEE est fortement encouragé.

La deuxième contribution de cette thèse est liée à la génération des circuits durcis. Un flux de conception au niveau transistor permet le développement de n'importe quelle structure ou méthodologie au niveau transistor. Différent des cellules standard conventionnelles, ces nouvelles techniques peuvent être directement appliquées dans le flux de conception.

On propose une nouvelle technique pour protéger les éléments séquentiels dans lesquels la redondance temporelle est insérée à l'intérieur des structures de bascules. On propose aussi une nouvelle méthodologie analytique de sizing de transistor afin de faire face aux SEEs dans les blocs combinatoires.

La principale caractéristique de ce nouveau modèle de sizing est la possibilité d'optimiser indépendamment les structures pull-up et pull-down de chaque porte. Ceci permet d'obtenir un circuit combinatoire durci aux rayonnements avec des pénalités réduites concernant la fonctionnalité de circuit.

A.2 Génération Automatique du Layout au Niveau Transistor

L'idée principale du travail proposé dans cette thèse est d'explorer les caractéristiques au niveau transistor afin de réduire la différence de performance entre les cellules standard et la conception *full custom*. Ainsi, on propose un flux de conception à partir d'une description RTL jusqu'au layout.

Le flux de conception inclut la génération d'une base de données utilisée dans la synthèse logique et dans la génération du layout. Pour ceci, le flux de conception est basé sur d'outils universitaires et commerciaux en visant la réduction de retard et de puissance.

Le *timing closure* est obtenue avec le sizing de transistor d'un circuit dans un ample nombre de possibilités. En d'autres termes, un flux de conception au niveau transistor permet d'optimiser les transistors dans les chemins critiques du circuit et maintenir les autres chemins avec des tailles minimum. La réduction sur la puissance est une conséquence directe de ce sizing optimisé.

L'optimisation du circuit est la première étape dans notre flux de conception. Ainsi, la description du circuit est pris avec transistors des tailles minimum. Pour chaque itération, le plus long chemin est extrait, des transistors sont dimensionnés et la puissance est analysée. Le résultat permet au concepteur de choisir le meilleur caractéristiques sur le retard et la puissance selon les spécifications du projet.

La prochaine étape est la génération du layout. Dans cet étape, les transistors sont placés et routés. L'extraction du layout permet l'évaluation du circuit et donne aux concepteurs importantes informations sur le processus d'optimisation. Si le retard et puissance sont d'accord avec les spécifications, le flux est fini. Autrement, les données extraites sont employées à une nouvelle phase d'optimisation.

A.2.1 La Génération de la *Super Bibliothèque*

Le layout des cellules n'existe pas dans le *superlib*. Autrement, seulement la fonction logique, les structures des cellules (connexions des transistors) et le retard et puissance estimée sont connues. Ces informations sont assez pour permettre la synthèse logique.

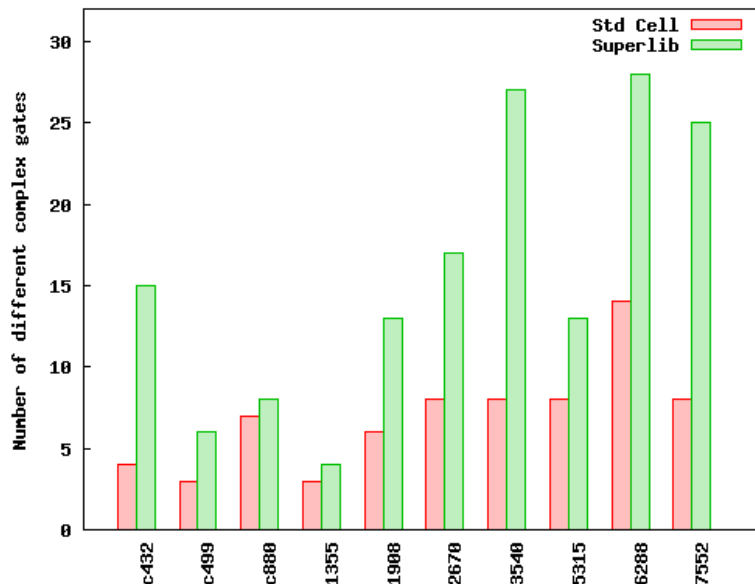


Figure A.3: Nombre de portes complexes différents dans les circuits pour une technologie commercial $0.35\mu m$.

La capacités des algorithmes de synthèse sur explorer un nombre large de fonctions logiques ne sont pas aussi claires dans la littérature. Une question importante au sujet de l'efficacité de notre *superlib* était si les outils commerciaux profit vraiment d'un nombre large de fonctions logiques au moment de la synthèse.

La figure A.3 prouve que les outils commerciaux peuvent explorer la synthèse quand un grand nombre des fonctions logique sont disponibles. Ces résultats ont été obtenus avec l'outil Cadence RTL Compiler [CAD07b]. Analysant ces dix circuits, nous remarquons que les circuits synthétisé avec notre *superlib* ont plus 52% de fonctions logiques. Des portes simples telles que des NAND, NORs et inverseurs ne sont pas incluses dans ces résultats.

La *superlib* contient les informations de retard sur différentes versions de chaque fonction logique. Cette bibliothèque est composé par 3.503 fonctions logiques différentes, qui est composé par chaque fonction logique avec jusqu'à quatre transistors empilés.

Le processus de génération de la bibliothèque (plus de 750.000 simulations) a été automatisé par des scripts. La génération de la *superlib* a pris autour 6 jours dans un SunfireV890 avec 8Gb de RAM.

A.2.2 Le flux de Projet au Niveau Transistor

L'optimisation au niveau transistor a été considérée dans plusieurs travaux [HO01, VWSS04, RBB05]. Vujkovic reporte dans [VWSS04] que la capacité par unité de

longueur change autour 35 fois dans une technologie $0.18\mu m$. Ceci montre comment l'optimisation de chaque cellule comme fonction de la capacité est importante, particulièrement dans le chemin critique.

Dans cette section nous présentons un flux complet de conception basé sur l'optimisation de transistors. Le but principal de cette méthodologie est d'explorer les possibilités d'un flux de conception au niveau transistor pour traiter les défis actuels en technologies submicroniques.

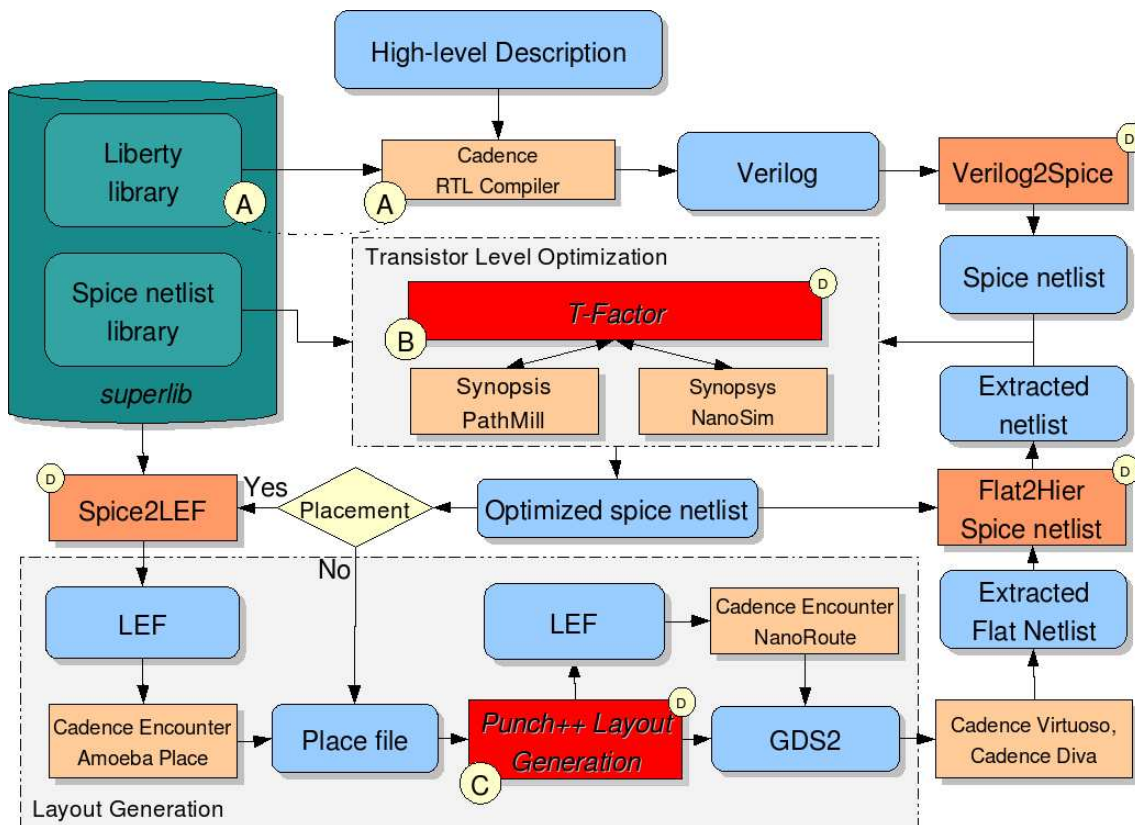


Figure A.4: Le flux de projet au niveau transistor.

Une fois que le *superlib* était créé, le layout du circuit peut être généré. La figure A.4 montre avec détails le flux de conception proposé dans ce travail. La méthodologie de conception proposée consiste sur un ensemble d'outils universitaire et commerciaux qui traitent défis submicroniques au même temps qui offre aux concepteurs un flux complet de conception.

Les principales différences entre le flux traditionnel de cellules et le flux au niveau transistor sont les suivantes :

- La possibilité de synthétiser un circuit avec un grand nombre de cellules en employant la *superlib* (figure A.4 A);
- L'outil d'optimisation de transistor appelé **T-Factor** capable de faire le sizing d'une circuit;

- Un nouveau outil de génération de layout appelé **Punch++** capable produire n'importe quelle type de porte statique CMOS complémentaire (Figure A.4 Label C).

La première étape dans le flux de conception est la génération de une description du circuit dans le format spice. Dans le flux proposé de conception, des circuits décrits dans des langues de haut niveau telles comme VHDL ou Verilog sont convertis en netlist au niveau logique en employant le Cadence RTL Compiler [CAD07b]. Après la synthèse, le netlist verilog est converti en description spice.

L'optimisation au niveau transistor commence par la description spice. L'outil **T-Facteur d'outil** les transistors dans les chemins critiques et ces transistors sont dimensionné afin de répondre à des caractéristiques de retard tandis que des transistors dans des chemins moins importants sont maintenus avec les tailles minimum.

Le placement de cellules est fait avec estimation de leur surface avec le Cadence amoeba placer [CAD07b]. Après le placement, le layout du circuit entier est produite et routé avec l'outil Cadence Nanoroute. La dernière étape du flux de conception est la compactage, le LVS et l'extraction de parasitics, qui rend possible d'évaluer l'exactitude du circuit.

A.2.2.1 L'Optimization au Niveau Transistor

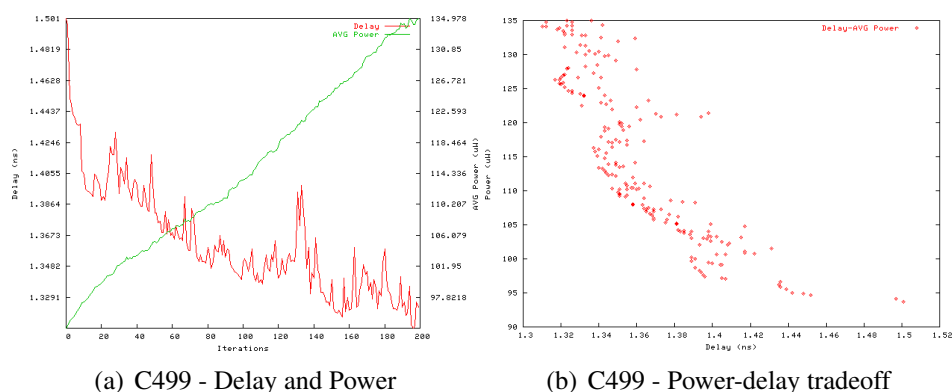


Figure A.5: Retard et puissance pour quelques circuits.

La figure A.5 montre la puissance et retard comme résultat de la première étape d'optimisation pour quelques circuits. Figure A.5(a) présente la puissance et le retard de chaque itération. Un aspect important est le linéaire développé de la puissance comme fonction de l'algorithme de sizing.

La quantité énorme de puissance dépensée est inévitable en recherchant des circuits plus rapides. Dans la plupart des conceptions, les plus petits retards ne sont pas les meilleures options due à la surface occupée et la puissance du circuit. En outre, quand on vise basse puissance des dispositifs, la conséquence est un plus grand retard du circuit.

Parfois le critère pour concevoir un circuit n'est pas la basse puissance dépensée ni une puissance exagérée, mais un bon rapport entre le retard et la puissance. Ainsi, figure A.5(b) présente un rapport entre le retard et la puissance dans lequel le concepteur peut

évaluer le processus d'optimisation. Chaque point dans le graphique consiste d'un possible circuit.

L'efficacité du processus d'optimisation au niveau transistor ne peut pas être mesurée si analysant seulement le plus mauvais retard d'un circuit. Autrement, l'efficace est prouvé en analysant le rapport parmi tous les chemins dans le circuit.

La fréquence d'opération dans un circuit est indiquée par le plus mauvais chemin retard. Un grand nombre de chemins avec très petit retardent peut signifier transistors trop dimensionné dû a inefficacité de l'algorithme de sizing. Une stratégie inefficace de sizing peut mener à une puissance excessive et indésirable.

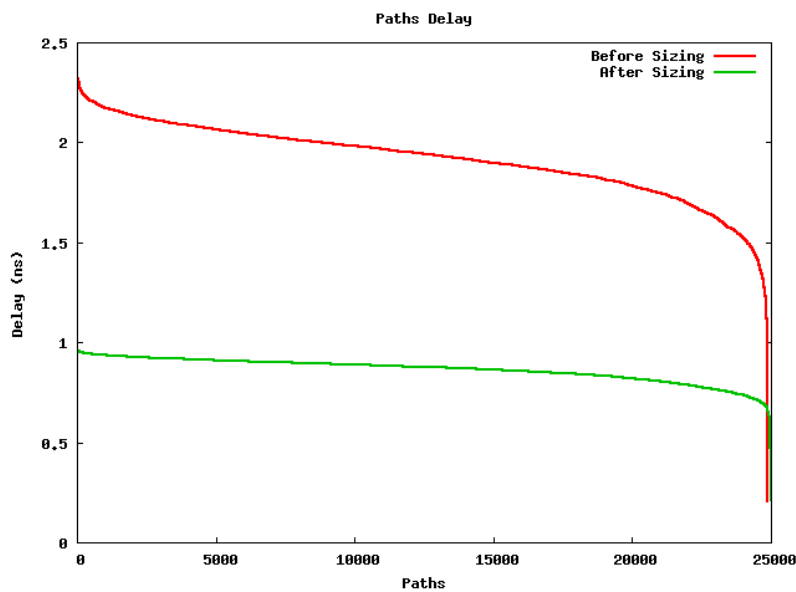


Figure A.6: Chemin de retard pour le circuit c1355.

La figure A.6 illustre le retard des chemins avant et après le sizing des transistors. L'axe X montre tous les chemins dans le circuit c1355 tandis que l'axe Y représente le retard de chaque chemin. La presque horizontale ligne donné par les retard des chemins après le sizing montre l'efficacité de l'algorithme. Cette uniformité dans les retard des chemins signifie que les portes dans les chemins les plus rapides présentent une taille proportionnée et ne consomment pas trop de puissance.

A.2.2.2 Optimisation au Niveau transistor pour la Réduction de la Courant de Fuite

La puissance statique est en train de devenir un importante facteur dans les circuit sub-microniques. La portion de dissipation de puissance de fuite a approché de la puissance dynamique, et les recherches estiment que la puissance statique excédera la puissance dynamique en technologie au-dessous de $65nm$ [KAB⁺03].

Pour ces raisons, la puissance statique doivent être incorporées dans la conception des systèmes dans le processus. Beaucoup de techniques ont été présentées dans les dernières années visant la basse de la puissance statique. Le problème de fuite de courant a été adressé à l'étape de conception par de diverses techniques telles que l'empilement de

transistor, l'utilisation de V_{DD} réduit et le contrôle de la longueur des transistors.

Le contrôle de la longueur de transistor consiste sur l'ajuste de la longueur des transistors pour réduire la fuite de puissance [GKSS04, KMS05, BCV06]. La fuite de courant est inversement proportionnelle avec la longueur du transistor. Cependant, le retard d'un transistor augmente avec la longueur du transistor. Ainsi, les chemins les plus rapides peuvent être dimensionné tandis que les transistors dans le chemin critique peuvent maintenir sa longueur afin de faire face au délai du circuit.

Bien que ces techniques soient possibles dans la conception des circuits, elles sont très complexes à employer. Pour la génération du layout au niveau transistor, ajuster la longueur des transistor est une méthode simple d'être utilisée dans le flux de conception parce que la génération du layout est la dernière étape du processus.

Pour cette raison, nous incluons l'ajuste de la longueur de transistor dans le flux de conception. Les transistors sont dimensionnés pour la réduction de la courant de fuite après que le circuit soit dimensionné pour le retard.

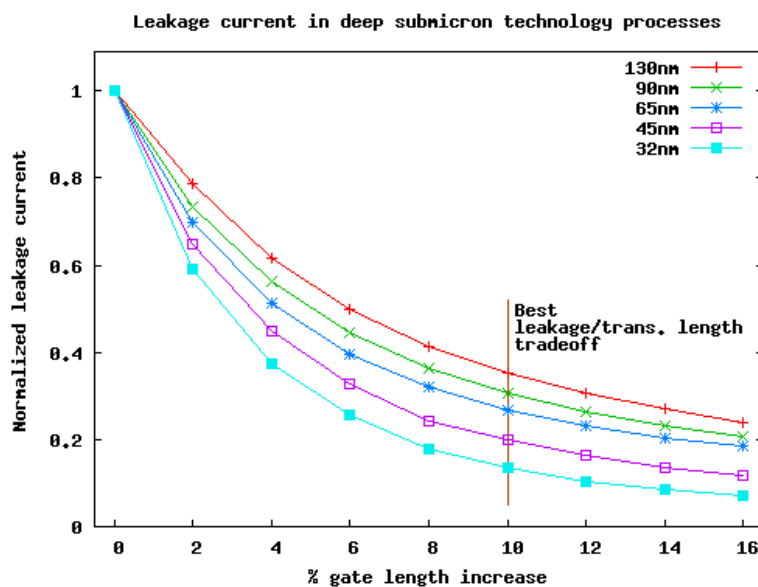


Figure A.7: Courant de fuite dans les circuits submicroniques.

La figure A.7 montre le courant normalisé de fuite dans les technologies submicroniques comme fonction de l'ajuste de longueur des transistors. Ces données ont été obtenues par des simulations spice avec les modèles prédictifs de technologie présentés par [ZC07].

Les résultats prouvent que l'ajuste de longueur des transistors est plus efficace avec l'évolution de technologie. Cependant, une limite supérieure au ajustement de longueur peut être défini à cause de la réduction exponentielle du courant de fuite.

Le gain dans la réduction de fuite est très petit quand la longueur de transistor est plus grande que 10% pour toutes les technologies. Pour cette raison, nous définissons une limite de 10% au ajustement de longueur des transistors pour nos expériences.

Résultats au sujet de l'ajuste de la longueur des transistors sont montrés dans le Tableau A.1. Les résultats montrent que l'ajuste de longueur des transistors est très ef-

Table A.1: La technique d'ajuste de la longueur des transistors pour quelques circuits ISCAS'85 avec la technologie 65nm.

Circuit	Cellules		Puissance		
	#Cell	(%)	Avant	Après	Réduction
C432	139/209	66%	3.5 μW	2.1 μW	60%
C499	208/296	70%	12.5 μW	8.1 μW	64%
C880	290/359	80%	10.6 μW	5.9 μW	55%
C1355	247/446	55%	10.3 μW	6.6 μW	64%
C1908	245/372	65%	14.4 μW	11.0 μW	76%
C3540	526/704	74%	20.1 μW	12.3 μW	61%
Réduction Moyenne					63%

ficace dans la réduction de la fuite de courant. Une moyenne de 70% des cellules étaient dimensionné et le circuit résultant dépensent 60% de la fuite de courant initiale. Il est important remarquer que le retard et la puissance dynamiques ne sont pas augmenté.

La distribution des chemins de retards avant et après l'ajuste de la longueur des transistors est montré dans la figure A.8. Les courbes montrées le nombre plus grand de chemins près de la cible de retard après l'ajuste de la longueur.

A.2.2.3 La Génération du Layout au Niveau Transistor

Comme précédemment discuté en ce chapitre, le layout du circuit est entièrement généré sur demande. Seulement le netlist avec la description des transistors et les règles de technologie sont nécessaires.

La génération du layout au niveau transistor consiste dans plusieurs étape comme illustré dans la figure A.4. D'abord, le layout ne peut pas être généré sans le placement des cellules. Ainsi, l'information estimée (retard, puissance et surface) de chaque cellule est employée pour placer les cellules dans le circuit. Afin d'estimer cette information, le netlist spice de chaque cellule est utilisé. La structure de transistors permet l'évaluation de retard et puissance. La surface occupé peut être estimé avec la structure de cellules et quelques règles de technologie. Cadence Amoeba [CAD07b] est utilisé pour placer les cellules.

Après le placement de cellules, la génération du layout peut être effectuée. Grands transistors ne sont pas faciles à placer et peuvent augmenter la surface occupé par les cellules. Pour ces raisons, la technique de *folding* est appliquée au netlist.

Le placement est appliqué au netlist où les cellules sont organisées dans lignes et placées côté a côté. Pour chaque ligne, les transistors sont placés et routé. Les transistors sont placés et routé sans pris en compte les règles de technologie. Seulement après le placement et la routage est que la compactage pris en compte les règles de technologie pour le layout,

Une fois que la génération de layout de chaque ligne est exécutée, le circuit est routé par complète dans le [CAD07b] et le circuit est écrit dans le format de GDSII. Le GDSII est le format standard employé par l'industrie.

Quand le layout est entièrement généré, le concepteur peut extraire les parasitics en

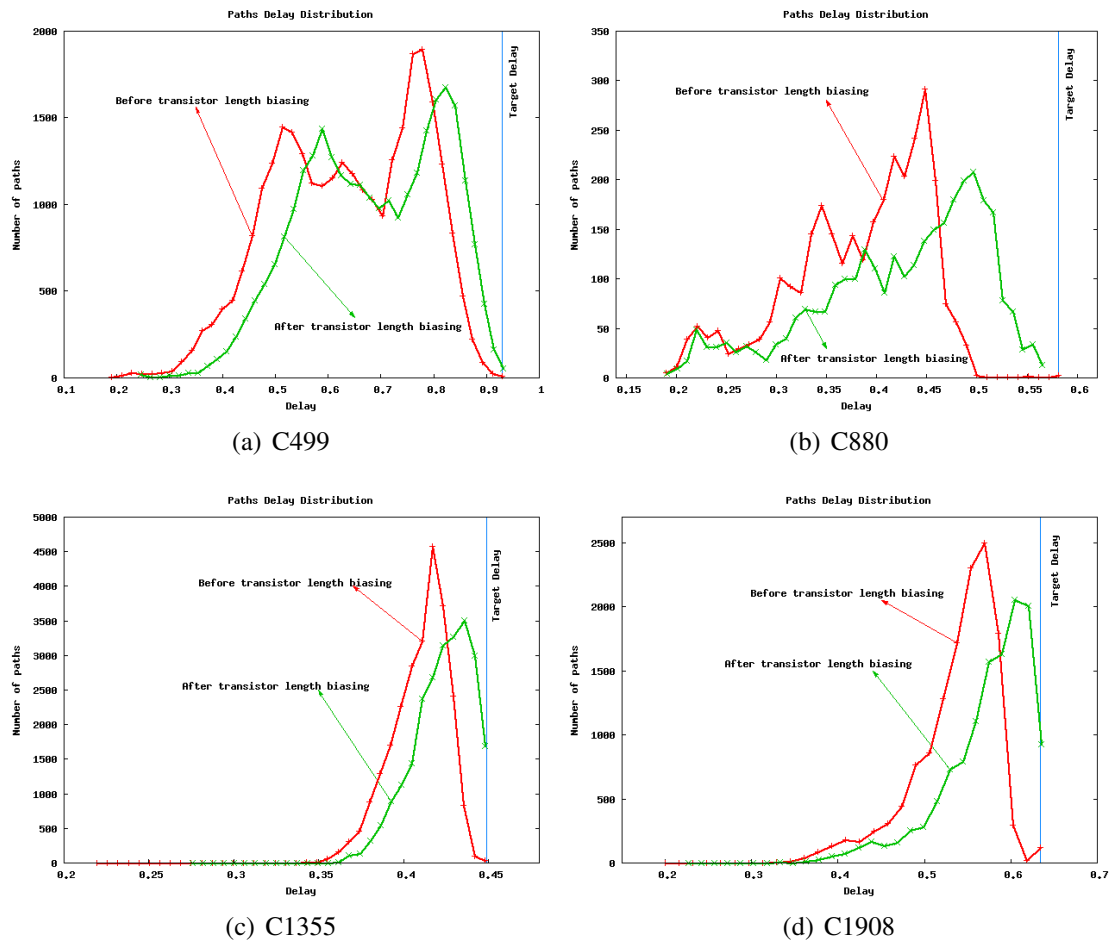


Figure A.8: Distribution des chemin de retard avant et après l'ajuste de la longueur de transistor.

utilisant l'outil Cadence DIVA. La validation du layout est également faite avec des outils de LVS et de DRC. Le netlist extrait est utilisé pour évaluer des caractéristiques électriques telles que la puissance et le retard. Si le circuit ne répond pas aux caractéristiques spécifiées, le concepteur peut répéter le processus d'optimisation afin d'améliorer les caractéristiques électriques du circuit.

A.2.2.4 Une comparaison entre le méthode traditionnel et le méthode proposé

Le Tableau A.2 présente quelques résultats sur le méthode proposée en comparaison avec l'approche standard de cellules pour une technologie commerciale de $0.35\mu m$. Les résultats sont très intéressants parce que montrent l'efficacité de notre flux de conception au niveau transistor.

Le processus de conception a été fait à basé sur l'effort élevé pour obtenir le retard minimum. Nous remarquons que cet effort élevé a eu comme conséquence l'insertion de beaucoup de *buffers* dans le chemin critique. Ceci explique le basse gain au sujet du retard (autour 11%) de notre méthodologie en comparaison avec l'approche de cellules standard. Le gain de puissance dans ces circuits est entre 15% et 42% à cause du nombre

Table A.2: Comparaison entre layout généré par approche de cellules standard et le flux de projet au niveau transistor

Circuit	Retard (ns)			Puissance Total (uW)		
	Std Cells	Proposé	Gain	Std Cells	Proposé	Gain
C432	3.97	3.68	7.3%	4416	3726	15.6%
C499	2.36	1.89	19.9%	11881	7122	40.0%
C880	1.88	1.85	1.5%	5592	3984	28.7%
C1355	2.50	2.45	2%	12071	6965	42.2%
C1908	2.39	2.06	13.8%	9493	6007	36.7%
C3540	5.15	4.05	21.4%	21141	15235	27.9%
C6288	9.46	7.98	15.6%	211593	145660	31.1%
Gain Moyenne			11.6%			31.7%

de portes complexes dans le circuit et la largeur optimisée de transistor.

A.3 Protection de Circuits Sequential aux SEEs

Plusieurs techniques ont été proposé pour la protection de circuits dans les dernières années. La plupart des techniques sont basées sur les structures redondantes. Cette redondance peut être temporelle ou spatiale.

La redondance *temporal* consiste en insérer une logique additionnelle dans la conception, qui garante l'évaluation d'un signal dans différents instants d'opération. Si une particule énergétique frappe un dispositif créant une impulsion de tension, cette logique additionnelle filtre les signaux et garantit l'atténuation de l'impulsion.

La redondance *spatial* est habituellement basée sur la réplique. L'idée principale dans la réplique spatiale est qu'une particule frappant un des éléments, n'affecte pas les autres. Ainsi, les sorties des éléments répliqués sont comparées et le signal filtré est propagé.

Toutes ces techniques sont caractérisées par la surface occupé élevé et important pénalités concernant retard et puissance. Les techniques sont habituellement très efficaces contre les SEE, mais les conséquences sont habituellement liées avec l'inefficacité du circuit concernant la fréquence et la puissance .

Une technique que profite de la nature temporelle des défauts transitoire est présenté en [Ang00]. Cette technique mene à une réduction significative de surface occupé comparée à la solution classique de TMR parce que l'idée principale est de combiner la conception des structures avec une redondance temporal.

Le fait que les erreurs affectent les sorties d'un circuit seulement pour une courte durée de temps peut être exploité en employant les éléments séquentiels asynchrones. Ces éléments produisent sur ses sorties un état déterminé pour chaque entrée correcte. Cet état correspond à l'opération mise au point de circuit. L'élément préserve son état précédent pour chaque entrée incorrecte.

Une manière de produire l'élément CWSP est de remplacer chaque transistor de la porte par une paire de transistors liés en série. Dans cette porte, quand les entrées d'une paire de transistors sont égales, les deux transistors se comportent comme un seule tran-

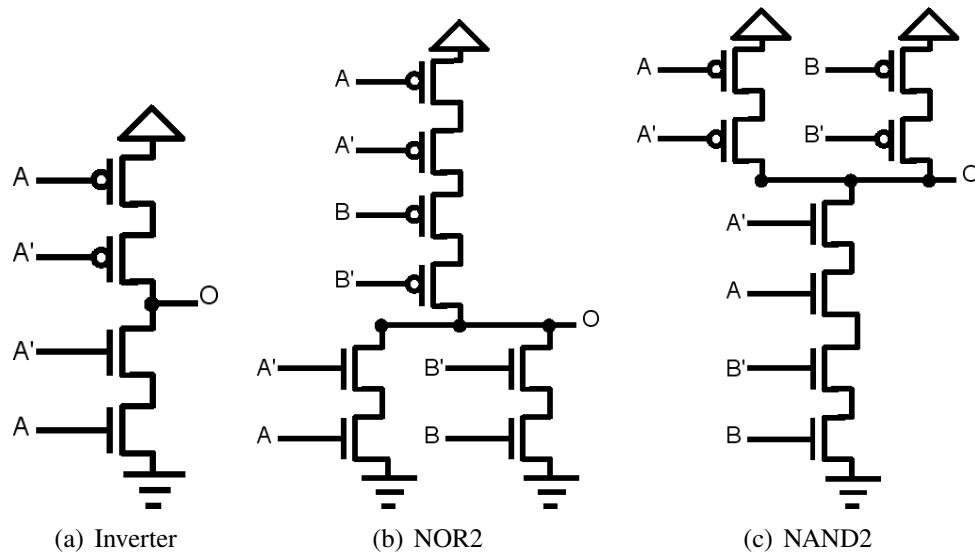


Figure A.9: Les portes INV, NOR2 and NAND2 avec la technique CWSP proposés par [Ang00].

sistor. Quand l'entrée d'un transistor est différente à cause d'une faute transitoire, l'autre transistor est bloqué et ne permet pas la propagation de la faute.

Figure A.9 montre quelques exemples. Dans l'approche de redondance de temps, au lieu de la duplication du circuit nous pouvons dupliquer le signal de sortie du circuit dans le domaine de temps, en observant ce signal à deux instants différents. Une des entrées de l'élément CWSP vient directement du circuit combinatoire produit tandis que l'autre entrée est retardée.

Table A.3: Surface occupée et retard pour les cellules standard INV, NAND and NOR en comparaison avec les cellules CWSP

	Surface (μm^2)			Retard (ps)		
	Std Cell.	CWSP	Over.	Std Cell	CWSP	Over.
INV	8.19	11.64	42%	83	160	92%
NAND	12.28	19.07	55%	98	172	75%
NOR	12.28	19.07	55%	102	260	154%

Les Tableaux A.3 présentent la surface occupée et le temps de propagation des cellules de CWSP accordant la figure A.9 avec une comparaison avec les cellules typiques d'une bibliothèque de cellules standard $0.18\mu m$. On montre que la surface occupée est entre 42% et 55% et le temps de propagation est entre 92% et 154% appliquant la même capacité aux deux portes. L'outil présenté dans [LDGR03, SWL⁺03, BLGR04] a été utilisé pour générer automatiquement ces cellules CWSP.

La génération des circuits combinatoires durcis avec la technique de redondance temporelle est présentée dans Figure A.10. Le bloc de retard doit pouvoir dégrader le signal à l'entrée de la cellule CWSP selon la période de la faute transitoire que nous désirons

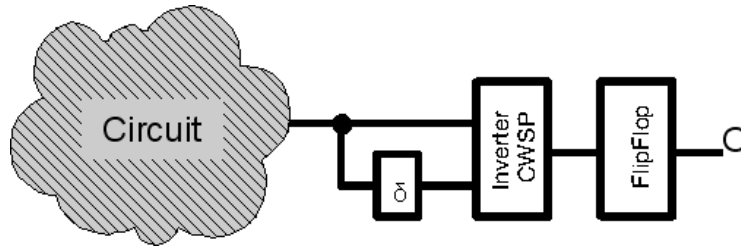


Figure A.10: Un exemple de circuit avec la redondance temporelle.

tolérer. La pénalité de temps est dans ce cas $D_{cw} + 2 \times D_{tr}$, où D_{cw} est le temps de transition et D_{tr} est la durée de la faute.

Table A.4: Surface et retard des cellules CWSP

	Trans. (ps)	Surface (μm^2)	Retard (ps)
INV	250	28.8	323
	500	46.0	538
NAND	250	59.2	370
	500	91.2	559
NOR	250	59.2	352
	500	91.2	572

Le Tableau A.4 présente la surface totale et le temps de propagation des nouvelles cellules CWSP développées comme représenté sur la figure A.10. Blocs de retard ont été développés et insérés dans les portes de CWSP afin d'obtenir les cellules durcis.

Nous avons proposé dans [LAR05a] une technique que vise la tolérance de perturbation aux logiques combinatoire et séquentielles. Elle emploie la technique de redondance temporelle présentée dans [Ang00] pour fournir la tolérance aux fautes. Pour tolérer les fautes transitoires dans les circuits combinatoires comme dans les éléments séquentiels, nous utilisons un latch modifié où le dernier inverseur est remplacé par un inverseur CWSP.

La figure A.11 montre la structure d'un latch en utilisant la logique de CWSP. La technique utilise un inverseur CWSP et blocs de retard pour réaliser la tolérance de fautes par la redondance temporelle.

Le Tableau A.5 présente la comparaison entre une Bascule classique trouvée dans une bibliothèque de cellules standard $0.18\mu m$ et une robuste proposée dans [LAR05a] (figure A.11). Nous supposons que le retard des blocs dans les Bascules TMR peuvent être partagés entre toutes les bascules dans le même domaine d'horloge, réduisant la surface occupée. Les résultats prouvent que les bascules robustes CWSP présentent un plus petit surface contre des fautes de $250ps$ en comparaison avec la technique TMR.

Une étude de cas a été faite afin de vérifier les pénalités de l'insertion des cellules CWSP dans un microprocesseur MIPS et un contrôleur 8051. La synthèse logique a été faite avec l'outil Synopsys Design Compiler [SYN07] et le layout a été généré par l'outil Cadence Silicon Ensemble [CAD07b].

Table A.5: Surface occupé par les bascules TMR et CWSP

	Surface	
	μm^2	Overhead
Standard Cell	57.6	—
CWSP (250ps)	181.7	215%
CWSP (500ps)	249.6	333%
TMR (250ps)	206.1	258%
TMR (500ps)	206.1	258%

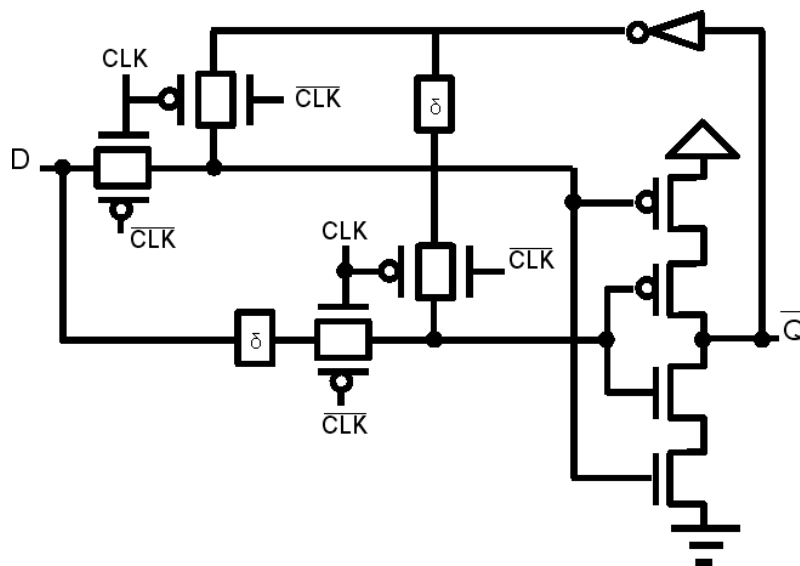


Figure A.11: An exemple de logique CWSP dans le latch comme proposé par [LAR05a].

Le Tableau A.6 montre la surface la fréquence pour les architectures MIPS et 8051 synthétisé dans une technologie $0.18\mu m$. La surface occupé par le circuit TMR est constant par le fautes de 250ps et 500ps parce que nous supposons que les blocs de retard sont partagés par toutes Bascules. Nous supposons également que trois lignes d'horloge ne sont pas un problème dans la conception de ces microprocesseurs.

Les résultats dans le Tableau A.6 montre que nous pouvons traiter le problème des fautes transitoire dans les blocs combinationnelles et séquentielles en utilisant la Bascule robuste CWSP avec des plus petites pénalités concernant surface et fréquence qu'en utilisant la technique de TMR.

La technique de TMR avec trois signaux d'horloge peut être un problème dans de plus grandes circuits. Ainsi, étapes additionnelles comme insertion de buffer ou sizing peut être nécessaires afin de garantir le fonctionnement de l'arbre d'horloge.

Cette étude de cas montrée la génération du layout des cellules durcis à utiliser dans la synthèse des circuits intégrés. L'importance d'un processus automatisé pour produire ces genre de cellules est liée au besoin de la production des circuits durcis pour plusieurs applications.

Table A.6: Les techniques CWSP et TMR dans les processeurs

	MIPS				8051			
No. El. Comb.	11,968				5,408			
No. bascules	1,793				1,359			
	Surface		Fréquence		Surface		Fréquence	
	μm^2	Over	MHz	Penalty	μm^2	Over	MHz	Penalty
Classic	480,317	—	77.7	—	234,720	—	58.2	—
CWSP (250ps)	746,480	55%	75.8	2.4%	436,240	85%	57.2	1.8%
CWSP (500ps)	890,172	85%	72.7	6.7%	550,560	134%	55.4	5.0%
TMR (250ps)	808,200	68%	73.9	5.1%	491,400	109%	56.0	3.8%
TMR (500ps)	808,200	68%	71.2	9.0%	491,400	109%	54.5	6.8%

A.4 Une Méthodologie Efficient de Sizing pour la Protection des Circuits

A.4.1 Combinational Circuits Sensitivity

L'analyse de sensibilité des circuits a été présentée dans plusieurs travaux [NJ06]. La plupart d'eux inclut la structure des portes et des détails du layout. L'analyse de la structure d'une porte consiste en évaluer dessus la propagation d'une faute comme fonctions des connexions des transistors. Par exemple, une faute dans le noeud de drain d'un transistor a plus de probabilité à être propagé qu'un noeud lié au V_{DD} ou GND .

Une analyse de sensibilité que considère le layout tient compte de la probabilité d'une particule frapper une région du layout. Par exemple, un grand surface de drain a une probabilité plus élevée qu'un plus petite.

Dans ce travail, la structure de la porte et son layout ne sont pas considérée. Différemment, nous considérons seulement le noeud de sortie de chaque cellule due à sa sensibilité plus élevée en comparaison des noeuds internes de la porte. Les caractéristiques du layout ne sont pas prises en considération dans l'analyse de sensibilité parce que nous considérons la sensibilité d'une porte après que le sizing devienne zéro comme fonction d'une charge critique donnée Q_c .

Il est important de remarquer que les aspects du layout ne sont pas considérés seulement pour l'analyse de sensibilité. Les détails du layout sont essentiels pour la méthodologie de sizing proposée.

Nous considérons la masquage logique et électrique comme la sensibilité d'un circuit. le masquage logique représente la probabilité d'une fautes transitoire être masqué par la fonction logique du circuit, et le masquage électrique décrit si une faute transitoire dans un noeud n'est pas propagée aux sorties primaires (PO). Ainsi, la sensibilité d'un circuit est donnée par

$$S_{circuit} = \sum_{n=1}^N (1 - L_n) \cdot (1 - E_n) \quad (A.1)$$

où L_n est le masquage logique et E_n est masquage électrique. L_n est une valeur de

Table A.7: Probabilité d'un noeud.

Fonction Logique	Probabilité
AND	$P_Z(1) = P_a(1) * P_b(1)$
NAND	$P_Z(1) = 1 - P_a(1) * P_b(1)$
OR	$P_Z(1) = 1 - (1 - P_a(1)) * (1 - P_b(1))$
NOR	$P_Z(1) = (1 - P_a(1)) * (1 - P_b(1))$
XOR	$P_Z(1) = P_a(1) + P_b(1) - 2 * P_a(1) * P_b(1)$
XNOR	$P_Z(1) = 1 - P_a(1) - P_b(1) + 2 * P_a(1) * P_b(1)$
BUF	$P_Z(1) = P_a(1)$
INV	$P_Z(1) = 1 - P_a(1)$

probabilité. Plus grand masquage logique signifie une plus petite probabilité d'une faute passagère être détectée dans les sorties du circuit. Le E_n est une valeur binaire où "0" indique que le transitoire est totalement atténuée et "1" indique que le faute peut être vue dans les sorties.

A.4.1.1 Le Masquage Logique

le masquage logique se produit quand un SET provoqué par une particule n'est pas propagé à une sortie primaire (PO) dû à la logique du circuit. La faute est masquée comme fonction d'un vecteur appliquée dans les entrées primaires (PI) du circuit. Techniques de contrôlabilité et d'observabilité sont employées pour définir le masquage logique d'un noeud.

La contrôlabilité dans les circuits de logique combinatoire dénote les capacités à un état soit placée dans un noeud. L'observabilité est une mesure pour quel point un état dans un noeud interne peut être connu aux sorties primaires (PO).

La contrôlabilité d'un noeud de sortie d'une porte est obtenue par la fonction logique des portes comme montré dans le Tableau A.7 [JJ05]. Ainsi, la propagation de la probabilité de contrôlabilité est faite pour le circuit entier.

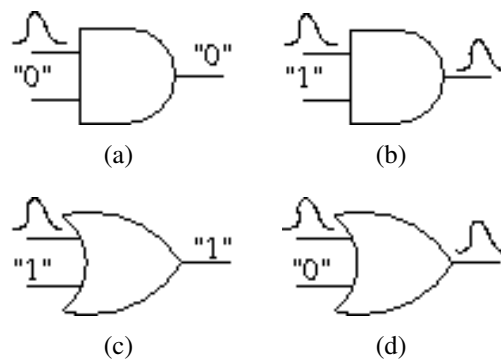


Figure A.12: Le masquage logique.

La figure A.12 illustre le masquage logique dans une porte. Une faute transitoire dans une des entrées de porte est propagée par la porte seulement si une valeur non contrôleur

est appliquée à l'autre entrée. La figure A.12(a) montre le masquage logique dans une porte AND comme fonction d'une valeur logique contrôleur "0" à l'entrée. Autrement, le masquage logique ne se produit pas si une valeur non contrôleur est appliquée (figure A.12(b)).

Dans la porte OR, la même situation est considérée, où la faute est propagée par la porte seulement si une valeur non contrôleur est appliquée à l'autre entrée. La figure A.12(c) montre que masquage logique comme fonction d'une valeur contrôleur et la Figure A.12(d) montre un cas où il n'y a aucun masquage logique.

A.4.1.2 Le Masquage Électrique

le masquage électrique est défini comme l'atténuation électrique d'une faute dans un noeud par les portes dans un chemin au point que le SET n'affecte pas les résultats du circuit.

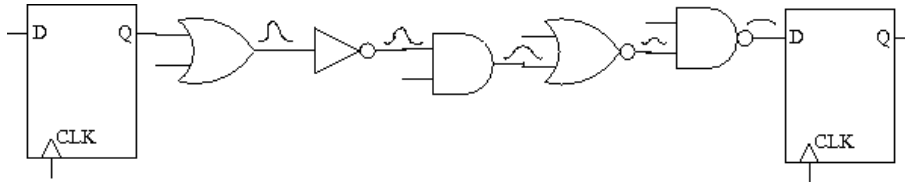


Figure A.13: Le masquage électrique.

La figure A.13 montre un exemple de dégradation de SET. Cette dégradation est la base de le masquage électrique, où la faute est dégradée comme fonction des caractéristiques électriques des portes dans le chemin. La faute peut être capturée par l'élément de mémoire si elle n'est pas assez dégradée.

A.4.2 Un modèle Analytique de SET

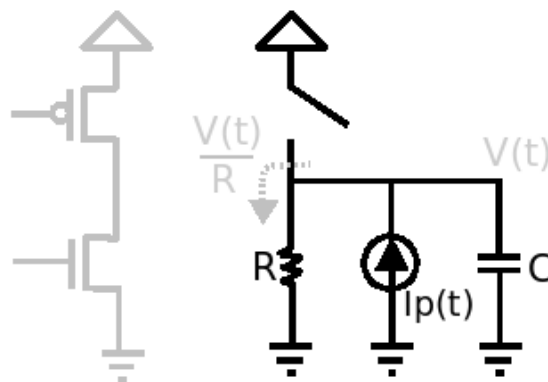


Figure A.14: Le circuit équivalent pour calculer la réponse d'une particule d'énergie.

Le modèle de sensibilité utilisé dans notre stratégie de sizing a été proposé par [WVK07]. Le modèle est basé sur deux paramètres de dispositif électrique. La capacité

C sur le noeud de sortie d'une porte g et la résistance R des transistor ouverts de cette porte.

La réponse d'un circuit comme fonction d'une particule d'énergie est modelée comme la réseau représenté dans la figure A.14, et peut être représentée par

$$\frac{V(t)}{R} - I_p(t) - C \frac{dV(t)}{dt} = 0 \quad (\text{A.2})$$

où $\frac{v(t)}{r}$ comprend le courant dérivé dans le transistor, représentés par la résistance R . $I_p(t)$ représente le courant provoqué par la particule frappant le dispositif et le dernier $C \frac{dV(t)}{dt}$ représente le courant dans le condensateur C .

La dérivation de cet modèle a une relation forte avec le comportement de dispositifs électriques et permet l'évaluation de la charge critique Q_c requise pour atténuer un SET dans un noeud.

A.4.2.1 Modelage de Résistances et Capacités

L'utilisation des résistances linéaires pour modeler des chemins de transistor est une méthode largement connue [WE93]. Ainsi, la résistance R peut être analytiquement déterminée par

$$R = \frac{1}{\mu_0 C_{ox} (\frac{W}{L}) (V_{gs} - V_{th})} \quad (\text{A.3})$$

où μ_0 est la mobilité du canal de transistor. C_{ox} est la capacité d'oxyde, qui est donnée par $\frac{\epsilon_0 \epsilon_{SiO_2}}{t_{ox}}$. ϵ_0 est la constante diélectrique, ϵ_{SiO_2} est la constante diélectrique relative d'oxyde et t_{ox} est l'épaisseur d'oxyde de porte. V_{gs} est la tension entre la porte et le source et V_{th} est la tension de seuil.

Tous ces paramètres sont des constantes reliées avec le technologie, excepté par le rapport de $(\frac{W}{L})$ qui représente les dimensions de transistor. Basé sur cet allongement, nous pouvons expliquer la relation entre la largeur de transistor et la résistance. Plus petit est la largeur d'un transistor, plus haute est la résistance.

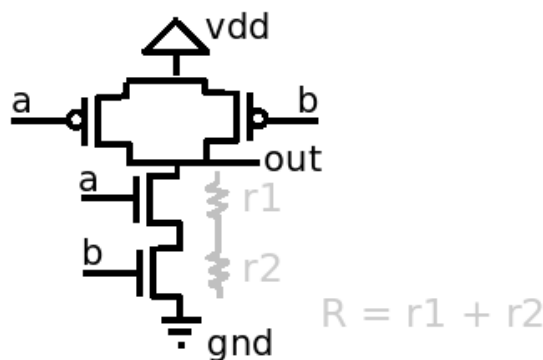


Figure A.15: Un transistor modelé par une résistance.

La figure A.15 illustre deux transistors empilés modelés comme résistances. Supposons que les transistors NMOS sont ouverts dans la porte de NAND dû aux signaux

Table A.8: Approximation de la capacité d'une porte MOS.

Parameter	Off	Non-saturated	Saturated
C_{gb}	$C_{ox}A$	0	0
C_{gs}	0	$\frac{1}{2}C_{ox}A$	$\frac{2}{3}C_{ox}A$
C_{gd}	$C_{ox}A$	$\frac{1}{2}C_{ox}A$	0
$C_g = C_{gb} + C_{gs} + C_{gd}$	$C_{ox}A$	$C_{ox}A$	$\frac{2}{3}C_{ox}A$

d'entrée $a = "1"$ et $b = "1"$. La résistance R est donnée par la somme des résistances r_1 et r_2 .

La capacité C est définies par la somme de trois capacités liées au noeud de sortie.

$$C = C_{diffusion_{g1}} + C_{connection} + C_{gate_{g2}} \quad (A.4)$$

où $C_{diffusion_{g1}}$ est la somme de toutes les capacités de jonction PN de la porte. $C_{connection}$ est la capacité de fil, et $C_{gate_{g2}}$ est la capacité de porte de tous les transistors liés au noeud de sortie.

Le $C_{diffusion_{g1}}$ est donné par

$$C_{diffusion_{g1}} = \sum_d^D \times C_{ja}A_d + C_{jp} \times P_d \quad (A.5)$$

où C_{ja} est la capacité de jonction par μ^2 , A_d est le surface de diffusion, C_{jp} est la capacité de périphérie par μ et P_d est le périmètre de diffusion.

Le troisième terme de la capacité C est la capacité de porte. Ainsi, $C_{gate_{g2}}$ est défini d'accord la région que la porte $g2$ fonctionne.

Le Tableau A.8 présente la capacité de porte d'accord la région d'opération. Basé dans cet information, la capacité de porte est définit par

$$C_{gate_{g2}} = \sum_{goff} C_{ox}A_g + \sum_{gon} \frac{2}{3}C_{ox}A_g \quad (A.6)$$

Ces équations analytiques permet de modéliser le comportement d'une faute transitoire comme fonction des caractéristiques électriques des dispositifs.

A.4.2.2 Le Modèle de SET

Messenger présente dans [Mes82] un modèle pour l'estimation SET. Le modèle représente les effets d'une particule α frappant un dispositif comme une double exponentielle courbe de courant. Cette courbe est obtenue par

$$I(t) = \frac{Q}{\tau_\alpha - \tau_\beta} \left(e^{-t/\tau_\alpha} - e^{-t/\tau_\beta} \right) \quad (A.7)$$

où Q est la charge injectée et peut être positif ou négatif, τ_α est la constante de temps de collection de la jonction et τ_β est la constant pour établir la voie d'un ion. τ_α et τ_β sont des constantes et dépendent de plusieurs facteurs d'une technologie.

Les Modèles présentés par [WVK07] sont des dérivations du double exponentiel pour obtenir le temps de crête t_{peak} et le crête de tension V_{peak} .

Il est important de remarquer que le τ_β est considéré comme beaucoup plus petit que le τ_α ($\tau_\alpha \gg \tau_\beta$) dans les formulations. En d'autres termes, le modèle assume un temps de montée très rapide au double exponentiel.

L'équation différentielle (A.2) est résolue pour obtenir la tension $V(t)$ dans le noeud. Ainsi, $V(t)$ est donné par

Le temps de crête t_{peak} dans le noeud a son valeur maximum dans

$$t_{peak} = \frac{\ln\left(\frac{\tau_\alpha}{RC}\right) \tau_\alpha RC}{\tau_\alpha - RC} \quad (A.8)$$

et, le crête de tension est obtenu quand on insère (A.2) dans le (A.8).

$$V_{peak} = \frac{I_0 \tau_\alpha R}{\tau_\alpha - RC} \left(\left(\frac{\tau_\alpha}{RC} \right)^{\frac{RC}{RC - \tau_\alpha}} - \left(\frac{\tau_\alpha}{RC} \right)^{\frac{\tau_\alpha}{RC - \tau_\alpha}} \right) \quad (A.9)$$

La charge critique Q_c est dérivée par (A.9) une fois que V_{peak} est connu. Thus, la charge critique Q_c est donné par

$$Q_c = \frac{V_{peak}(\tau_\alpha - RC)}{R \left(\left(\frac{\tau_\alpha}{RC} \right)^{\frac{RC}{RC - \tau_\alpha}} - \left(\frac{\tau_\alpha}{RC} \right)^{\frac{\tau_\alpha}{RC - \tau_\alpha}} \right)} \quad (A.10)$$

La tension au noeud frappé montre un comportement double exponentiel dans lequel la tension transitoire V_{peak} est atteinte au temps t_{peak} . La tension commence à diminuer exponentiellement après t_{peak} .

$$\tau_n = t_{peak} - RC \ln \left(\frac{\frac{1}{2} VDD}{V_{peak}} \right) - \tau_\alpha \ln \left(\frac{\frac{1}{2} VDD}{V_{peak}} \right) \quad (A.11)$$

L'équation (A.11) montre la durée de la faute transitoire τ_n , où le deuxième terme correspond à la solution analytique si le temps de RC est beaucoup plus grand que τ_α et le dernier terme correspond à la solution analytique si le temps τ_α est beaucoup plus grand que RC .

A.4.2.3 Single Event Transient Propagation

L'analyse de la propagation d'une faute transitoire prouve que la dégradation d'impulsion est directement influencée par le retard de propagation τ_g d'une porte. En d'autres termes, un plus grand τ_g mène à une plus grande dégradation de la faute.

Wirth *et al* ont proposé un modèle de dégradation d'impulsion basé sur l'ajustement de courbe [WVKN07]. Le modèle considère un paramètre k égal au minimum rapport τ_n/τ_g nécessaire pour propager le SET pour la prochaine étape dans un chemin du circuit.

La dégradation d'une faute transitoire est donnée par

$$\tau_{n+1} = \begin{cases} 0 & \text{if } (\tau_n \leq k\tau_g), \\ (k+1)\tau_g(1 - e^{-(\tau_n/\tau_g)}) & \text{if } ((k+1)\tau_g < \tau_n \leq (k+3)\tau_g), \\ \frac{\tau_n^2 - \tau_g^2}{\tau_n} & \text{if } ((k+1)\tau_g < \tau_n \leq (k+3)\tau_g), \\ \tau_n & \text{if } (\tau_n > (k+3)\tau_g). \end{cases} \quad (A.12)$$

Pour une transitoire avec une petite durée, la crête de tension est plut petite que $\frac{1}{2}v_{dd}$ et l'atténuation complète doit être considérée. Ainsi, le premier cas modèle les situations où le transitoire est totalement supprimée.

Le deuxième et troisième cas sont liés à une dégradation partielle dans le transitoire selon la relation entre la durée du SET τ_n et le retard de la porte τ_g . Le quatrième cas de dégradation consiste des situations où le transitoire n'est pas dégradée d'une étape à l'autre ou peut être négligé.

Ces quatre cas de dégradations sont la base à l'algorithme de sizing en raison de ses propriétés de propagation. Ces propriétés peuvent être utiles également pour obtenir la durée maximum acceptable d'une faute dans un noeud.

Une remarque importante est qu'une faute transitoire n'a pas besoin d'être atténuée dans le noeud n (à moins que dans les cas où des portes sont liées aux sorties). Le SET peut être atténué dans le chemin entier entre le noeud n et les sorties primaires.

A.4.3 The Transistor Sizing Strategy

La stratégie de sizing de transistor proposée dans cette thèse consiste de trouver les plus petites largeurs de transistor de chaque porte de circuit pour l'atténuation d'un SET. Les formulations précédemment discutées sont la base pour l'algorithme de sizing.

Algorithm 9 Le sizing de transistor pour l'atténuation de SET.

Require: Les portes G , Les nets N , Les sorties O , La sensibilité maximum M , La charge critique Q_c , La sensibilité désirée du circuit $S_{desired}$

Ensure: Les portes avec les transistors dimensionnés G_{new}

```

1:  $G_{new} \leftarrow \emptyset$ 
2: for all  $n \in N$  do
3:    $L_n \leftarrow \text{calculateLogicalMasking}(n);$ 
4:    $E_n \leftarrow \text{calculateElectricalMasking}(n, Q_c);$ 
5:    $S_n \leftarrow (1 - L_n) \cdot (1 - E_n)$ 
6: end for
7:  $V \leftarrow O$    {Nets to visit, starting from the outputs.}
8: while  $V \neq \emptyset$  do
9:   for all  $n \in V$  do
10:     $g \leftarrow \text{getFaninGateConnectedToNet}(n);$ 
11:    if  $S_n > M$  then
12:       $\tau_n \leftarrow \text{getMaximumSET}(n, g);$ 
13:       $g_{new} \leftarrow \text{sizeTransistors}(s, g, \tau_n);$ 
14:       $G_{new} \leftarrow G \cup \{g_{new}\} \setminus \{g\}$ 
15:    end if
16:     $I \leftarrow \text{getGateInputs}(g);$ 
17:     $V \leftarrow V \cup I \setminus \{n\}$ 
18:   end for
19: end while

```

La stratégie de sizing proposée est présentée dans l'algorithme 9. Les premières lignes (2-6) définissent la sensibilité du circuit comme montré dans (A.1). La stratégie de sizing

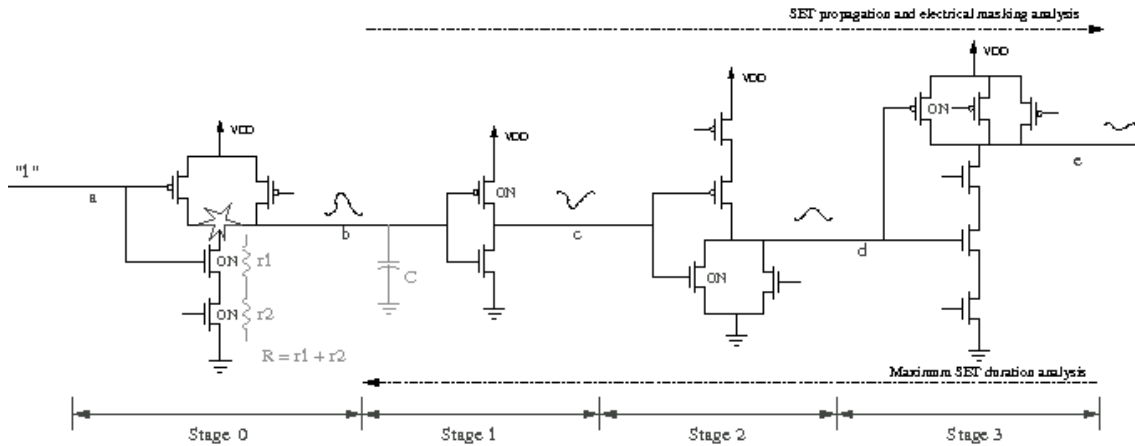


Figure A.16: Une exemple de propagation de transitoire.

de transistor commence à la ligne 8, où chaque noeud n du circuit est visité afin de trouver la largeur minimum de transistor pour chaque porte g liée à ce noeud. Il est important de remarquer que seulement les noeuds avec la sensibilité plus grande que la sensibilité maximum définie M sont évalués (ligne 11).

La Fonction `getMaximumSET(  $n$ ,  $g$  )` (ligne 12) trouve la durée maximum d'un transitoire τ_n dans le noeud n qui est supprimé avant les sorties primaires.

La Fonction `sizeTransistors(  $s$ ,  $g$ ,  $\tau_n$  )` (ligne 13) augmente la largeur de transistors jusqu'à le SET du noeud n soit plus petite que τ_n . Quand cette situation est atteinte, nous considérons que les transistors de la porte g sommes dimensionné comme prévu à la charge Q_c .

Les autres lignes de la stratégie montrée dans l'algorithme 9 donne une certaine idée au sujet de la navigation dans les noeuds. L'algorithme évalue chaque noeud de la logique combinatoire, des sorties primaires (PO) aux entrées primaires (PI). Ceci est fait parce que le retard des portes est changé après le sizing. Quand les transistors d'une porte sont dimensionnés, le retard devient habituellement plus petit et la propagation des transitoires a une plus petite dégradation.

L'interprétation incorrecte au sujet de la propagation du SET se produit si le transitoire est évaluée avant le sizing des portes. Ainsi, quand le SET est évalué dans un noeud n , nous garantissons que chaque porte dans le chemin entre ce noeud n et les sorties, ont été déjà dimensionnés.

A.4.4 Le Modèle de Sizing

La technique de sizing de transistor proposée est fondamentalement séparée dans trois étapes. Ces étapes sont liées à la sensibilité d'un noeud du circuit comme fonction d'une particule frappant le circuit et la largeur minimum de transistor requis pour atténuer un SET (masquage électrique).

La première étape est l'analyse de la sensibilité, qui est représentée dans les lignes 2 à 6 (algorithme 9). La sensibilité d'un noeud n est donnée par le masquage logique et électrique.

L'exemple dans la figure A.16 montre une particule avec charge Q que frappe la sortie d'une porte NAND au *stage 0*. Le transitoire est propagée du noeud b par chaque étape jusqu'à la sortie du circuit. La durée du transitoire au noeud frappé est définie par des équations décrites en sous-section A.4.2.2 et est fonction de la résistance R , de la capacité C , de la charge Q et de quelques constantes de processus de technologie.

La dégradation d'un SET dépend du retard de chaque porte dans le chemin comme expliqué dans la section A.4.2.3. Le masquage électrique est fondamentalement l'analyse de la dégradation du SET par le chemin. On assume qu'un noeud est électriquement masqué si le transitoire est supprimée avant les sorties.

Pour cette raison, la Fonction `getMaximumSET(  $n$ ,  $g$  )` (ligne 18) trouve la durée maximum d'un SET τ_n pour un filet n qui est atténué juste avant les sorties primaires. on suppose que la durée du SET aux sorties primaires $\tau_{out} = 0$. L'équation de la section A.4.2.3 ont été modifiées pour trouver une durée acceptable d'un SET dans le net n .

Nous avons dérivé ces équations pour obtenir la durée τ_n d'un SET aux entrées de chaque porte comme fonction de la durée τ_{n+1} du SET. Les cas où le transitoire est totalement atténuée ou propagée sans n'importe quelle dégradation sont montrés dans les premiers et derniers cas de (A.12). La dégradation partielle dans l'impulsion passagère est présentée dans les autres équations. Ces équations ont été dérivées comme suit.

$$\tau_n = \tau_g \left(k - \ln \left(1 - \frac{\tau_{n+1}}{\tau_g(k+1)} \right) \right) \quad (\text{A.13})$$

Les situations où τ_n est plus grand que $k\tau_g$ et plus petit que $(k+1)\tau_g$ sont traitées par (A.13).

$$\tau_n = \frac{\tau_{n+1} + \sqrt{\tau_{n+1}^2 + 4\tau_g^2}}{2} \quad (\text{A.14})$$

Les cas de propagation où τ_n est plus grand que $(k+1)\tau_g$ et plus petit que $(k+3)\tau_g$ sont traitées par (A.14).

La Fonction `sizeTransistors(  $s$ ,  $g$ ,  $\tau_n$  )` trouve la plus petite largeur de transistors pour une porte g comme fonction de la durée d'un transitoire τ_n . Dans la figure A.16, les transistors NMOS de la première porte sont ouverts et ces transistors sont les responsables principaux par la durée du SET. Ainsi, seulement les transistors NMOS sont dimensionnés afin de réduire les résistances $r1$ et $r2$ (la capacité est indirectement augmentation comme fonction de la surface des diffusions).

Le sizing des transistors est modelé comme basé a les équations présentées en section A.4.2. L'algorithme consiste sur appliquer le méthode de la bisection pour trouver la largeur de chaque transistor pull-up et pull-down de la porte.

A.4.5 Results

Les Tableaux A.9 et A.10 montrent les résultats obtenues par la stratégie de sizing proposée. Les résultats incluent une comparaison entre les méthodologies symétriques et asymétriques de sizing pour un processus de technologie 180nm [ZC07]. Le paramètre de

Table A.9: Le méthode de sizing proposé pour l'atténuation de SET avec charge critique $Q = 0.3pC$ et circuit avec sensibilité de 50%.

Combinational Circuit	Sizing Methodology	Overhead		
		Area (%)	Power (%)	Timing (%)
C432	Symmetric	47.4	63.8	0.0
	Asymmetric	35.5	50.7	2.0
C880	Symmetric	88.0	72.4	0.0
	Asymmetric	69.2	51.6	0.0
C1355	Symmetric	62.4	38.6	16.0
	Asymmetric	50.6	29.5	15.8
C1908	Symmetric	47.0	35.5	12.0
	Asymmetric	37.0	29.0	8.8
Average overhead	Symmetric	61.2	52.7	7.0
	Asymmetric	48.0	40.2	6.65

propagation de transitoire k a été défini par simulations hspice comme 0.8 pour cette technologie. Le sizing des transistors était viser la réduction de la sensibilité à 50% (Tableau A.9) et à 0% (Tableau A.10).

Comme discuté dans la section 4.2.1.2, une étude présentée dans [ZM06] prouve que la charge déposée très de peu de particules est plus élevée que $0.3pC$ au niveau de l'atmosphère. Nous employons cette valeur dans nos expériences en considérant comme le plus mauvais cas concernant la charge déposée.

Le premier point important montré par ces résultats concerne la basse surface présentée par la méthodologie proposée. Le plus mauvais cas était de 87% pour la protection complète (sensibilité de 0%) contre des particules avec la charge $Q = 0.3pC$. Les résultats montrent une surface occupée moyen de 83% pour le sizing asymétrique et le 61% pour le sizing symétrique. La puissance présente à 70% pour le sizing symétrique contre 43% pour l'asymétrique. Résultats montrés aussi une très de petites pénalité de retard de 10% pour le circuit avec sensibilité de 0%.

Le sizing asymétrique de transistor a eu comme conséquence un occupation de surface, consommation de puissance et retard plus petit en comparaison du sizing symétrique. Ces résultats montrés l'efficacité du sizing asymétrique.

A.5 Conclusion

Les contributions de cette thèse sont fondamentalement divisées dans deux majeures parties. La première est lié à l'élaboration d'une nouvelle méthodologie capable de produire circuits intégrés optimisés concernant le retard et la puissance. Le flux de conception au niveau transistor optimise chaque porte dans le circuit d'accord les capacités dans lesquelles elle est impliqué.

L'arrivée des technologies submicroniques profondes a inclus plusieurs défis dans la conception des circuits. Les géométries sont réduites, les alimentations d'énergie deviennent plus petites et la densité logique atteint un taux très élevé. Le plus grand nombre de

Table A.10: Le méthode de sizing proposé pour l'atténuation de SET avec charge critique $Q = 0.3pC$ et circuit avec sensibilité de 0%.

Combinational Circuit	Sizing Methodology	Overhead		
		Area (%)	Power (%)	Timing (%)
C432	Symmetric	69.8	105	1.2
	Asymnetric	50	59.7	0.0
C880	Symmetric	115.3	88.7	12.3
	Asymmetric	86.9	59.1	13.2
C1355	Symmetric	80.0	61.6	24.8
	Asymmetric	58.6	37.2	17.1
C1908	Symmetric	69.2	20.89	13.0
	Asymmetric	49.2	17.4	10.16
Average overhead	Symmetric	83.5	69.0	12.82
	Asymmetric	61.1	43.3	10.11

couche de métal, associé à ces caractéristiques submicroniques, décalent le paradigme de conception des circuits avec retard dominé par la logique pour le retard dominé par des interconnexions.

Quelques travaux récents prouvent que la capacité par unité de longueur peut varié aléatoirement jusqu'à 35 fois dans une technologie $180nm$ [VWSS04]. Cet énorme variation souligne le besoin d'optimisation au niveau transistor parce qu'il est pratiquement impossible de prévoir ces conditions avant la phase du layout.

Le flux de conception proposé est basé sur des outils universitaires et commerciaux. Les outils commerciaux incluent la synthèse de logique, le placement et routage, et l'analyse de retard et puissance. Les outils universitaire ont été développés pour faire face aux lacunes entre la méthodologie de génération conventionnelle et la méthodologie de génération au niveau transistor.

Fondamentalement, le flux de conception au niveau transistor présente trois différences en comparaison avec la méthodologie de cellules standard:

1. *La génération de la bibliothèque:* Le layout des cellules n'est pas produite au temps de génération de la bibliothèque. Ceci permet de manière significative d'augmenter le nombre de fonctions logiques.
2. *Optimisation des transistors:* L'optimisation au niveau transistor permet de trouver la largeur optimisée des transistors concernant les capacités impliquées à la porte.
3. *Génération du layout:* La génération du layout est effectuée après l'optimisation des transistors pour faire face à une gamme très élevés des possibilités au sujet de la largeur de transistors.

On propose également une réduction de la courant de fuite avec une technique que permet l'ajuste de la longueur des transistors.

Tous ces dispositifs permet d'atténuer les problèmes de retard et puissance dans les circuits DSM. Les résultats prouvent que cette méthodologie est très prometteuse. Les

comparaisons entre la méthodologie au niveau transistor et des cellules standards montré quelques résultats intéressants où la méthodologie proposée présente autour 11% d'amélioration concernant le retard du circuit et plus de 30% de réduction de puissance.

La deuxième contribution de cette thèse concerne l'application du flux de conception au niveau transistor dans la protection des circuits intégrés contre les SEE. L'aspect principal de cette possibilité est l'élaboration d'une nouvelle méthodologie de sizing de transistor pour produire des circuits combinationnels durcis.

L'évolution des technologies a également des effets dans des circuits intégrés au sujet de l'échec fonctionnel dû aux SEEs. La réduction de la longueur de porte et les basses tensions d'alimentation rendent les circuits sensibles aux particules énergiques qui était sans valeur dans les technologies plus anciennes.

Deux contributions sont présentées au sujet de la protection contre SEE. La première est lié à l'insertion de la redondance temporel dans les éléments séquentiels. L'idée principale consiste d'appliquer les concepts proposés par Anghel au sujet de la technique CWSP dans les bascules [Ang00].

Cette méthodologie présente basse surface occupé, mais les pénalités de retard sont totalement dépendent de la durée du transitoire que nous voulons atténuer.

La deuxième contribution sur la protection de SEE implique dans une nouvelle méthodologie de sizing de transistor. La sensibilité des circuits combinationnels sont définies par l'analyse de masquage logique et électrique. Le masquage logique donne la probabilité d'un transitoire être masqué par la logique de circuit, qui est calculée par des techniques de contrôlabilité et d'observabilité. Le masquage électrique décrit si une faute transitoire dans un noeud n'est pas propagée aux sorties primaires dû à l'atténuation électrique.

Le méthode de sizing est basée sur le modèle analytique de détection de faute transitoire présentée par [WVK07, WVNK07]. Fondamentalement, le modèle est basé sur deux paramètres de dispositif électrique. La capacité C dans le noeud de sortie d'une porte g et la résistance R des transistors ouverts de cette porte.

Une méthode de propagation est employée, dans lequel le retard de la porte et la durée d'un transitoire sont évaluées. Le modèle de propagation est très important parce qu'il considère que le SET doit être atténuée pendant le chemin entier et pas dans le noeud frappé par la particule. Ainsi, des transistors excessivement dimensionnés sont évités. Le modèle considère également les blocs de transistors pull-up et pull-down indépendamment. Seulement les transistors directement liés à l'atténuation du SET sont dimensionné.

Les résultats montrent la surface occupé, le retard et la puissance en comparaison d'une méthodologie symétrique dans lequel permettant le développement des circuits à haute fréquence. Ces deux contributions principales sont la base de cette thèse, mais d'autres sujets sont présentées et discutées avec détails.

RESUME EN FRANÇAIS:

Les technologies submicroniques ont inséré des nouveaux défis dans le projet de circuits intégrés à cause de la réduction des géométries, la réduction de la tension d'alimentation, l'augmentation de la fréquence et la densité élevée de la logique. Un des buts de cette thèse est liée au développement des outils d'EDA capables de faire face à ces nouveaux défis. Cette thèse est divisée dans deux contributions principales. La première contribution est liée à l'élaboration d'une nouvelle méthodologie capable de produire des circuits optimisés en ce qui concerne le retard et la puissance. On propose un nouvel flow de conception dans lequel le circuit est optimisé au niveau transistor. La deuxième contribution de cette thèse est reliée avec le développement des techniques pour les circuits durcis aux rayonnements. La technique *Code Word State Preserving* (CWSP) est utilisé pour appliquer la redondance dans les bascules. On propose aussi une nouvelle méthodologie dans lequel la taille de transistor est dimensionné pour l'atténuation de faute type *Single Event Transient*. La méthode de *sizing* est basée sur un modèle analytique. Le modèle considère indépendamment et les transistors PMOS et NMOS des portes.

TITRE EN ANGLAIS:

Transistor Level Automatic Generation of Radiation-Hardened Circuits

RESUME EN ANGLAIS:

Deep submicron technologies have increased the challenges in circuit designs due to geometry shrinking, power supply reduction, frequency increasing and high logic density. One of the goals of this thesis is to develop EDA tools able to cope with these DSM challenges. This thesis is divided in two major contributions. The first contribution is related to the development of a new methodology able to generate optimized circuits in respect to timing and power consumption. A new design flow is proposed in which the circuit is optimized at transistor level. The second contribution of this thesis is related with the development of techniques for radiation-hardened circuits. The Code Word State Preserving technique is used to apply timing redundancy into latches and flipflops. Further, a new transistor sizing methodology for Single Event Transient attenuation is proposed. The sizing method is based on an analytic model. The model considers independently pull-up and pull-down blocks.

SPECIALITE:

Micro et nano electronique.

MOTS-CLES:

Génération de circuits intégrés, tolerance aux fautes, fautes transitoires, technologies sousmicroniques.

LABORATOIRE:

Laboratoire TIMA
46, Avenue Félix Viallet
38031 Grenoble Cedex, France

ISBN: 978-2-84813-113-9

ISBN: 978-2-84813-113-9