



Design, fabrication and characterization of a hybrid III-V on silicon transmitter for high-speed communications.

Thomas Ferrotti

► To cite this version:

Thomas Ferrotti. Design, fabrication and characterization of a hybrid III-V on silicon transmitter for high-speed communications.. Other. Université de Lyon, 2016. English. NNT : 2016LYSEC054 . tel-01529424

HAL Id: tel-01529424

<https://theses.hal.science/tel-01529424>

Submitted on 30 May 2017

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.



ÉCOLE
CENTRALE LYON

N° d'ordre NNT : 2016LYSEC54

THESE de DOCTORAT DE L'UNIVERSITE DE LYON
opérée au sein de l'Ecole centrale de Lyon

Ecole doctorale N°160
Electronique, Electrotechnique, Automatique (EEA)

Spécialité de doctorat : Electronique, micro et nano-électronique, optique et laser

Soutenue publiquement le 16/12/2016, par :

Thomas Ferrotti

**Design, fabrication and characterization of a hybrid III-V on
silicon transmitter for high-speed communications**

Devant le jury composé de :

Marris-Morini, Delphine	Professeur	C2N	Rapporteur
Morthier, Geert	Professeur	Ghent University - IMEC	Rapporteur
Grandjean, Nicolas	Professeur	EPFL	Président du jury
Liang, Di	Ingénieur de recherche	HPE	Examineur
Oxenløwe, Leif Katsuo	Professeur	DTU	Examineur
Seassal, Christian	Directeur de recherche	CNRS - INL	Directeur de thèse
Ben Bakir, Badhise	Ingénieur – Chercheur	CEA-LETI	Encadrant de thèse
Bœuf, Frédéric	Ingénieur de recherche	STMicroelectronics	Invité



Contents

Contents.....	iii
Thanks/Remerciements	vii
Chapter I: Introduction	1
I-1. Context	2
I-1.1. Information society and photonics	2
I-1.2. Current status of the photonics market	2
I-2. Silicon photonics	3
I-2.1. Why should we use silicon for photonics?	3
I-2.2. Silicon photonics device library	4
I-2.3. Optical source issue in silicon photonics	5
I-3. Heterogeneously integrated III-V on silicon laser	9
I-3.1. General principles of laser operation	9
I-3.2. Hybrid III-V on silicon laser structure description	13
I-3.3. Existing hybrid III-V on silicon laser cavity designs and state-of-the-art.....	17
I-3.4. Advantages of hybrid III-V on silicon lasers and their integration issues	21
I-4. High-speed integrated transmitters	22
I-4.1. State-of-the-art for integrated transmitters in silicon photonics	22
I-4.2. Work objectives and targeted specifications.....	23
Chapter II: Silicon Mach-Zehnder modulator	25
II-1. Design principles and state-of-the-art.....	26
II-1.1. General principles of optical modulation	26
II-1.2. Figures of merit.....	27
II-1.3. Physical phenomena for optical modulation in silicon	29
II-1.4. Device structures used for carrier depletion devices	33
II-1.5. Target performance of the modulator	35
II-2. Silicon Mach-Zehnder modulator design	35
II-2.1. Junction design	35
II-2.2. Speed limitations and travelling-wave electrodes design	39
II-2.3. Device architecture.....	46
II-3. Silicon Mach-Zehnder modulator fabrication	47
II-4. Mach-Zehnder modulator characterization.....	51
II-4.1. Passive components optical characterization	51
II-4.2. Active region static electro-optical characterization	52

II-4.3. Small-signal electrical and electro-optical characterization	53
II-4.4. Large-signal electro-optical characterization	56
II-5. Conclusions and possible improvements	62
Chapter III: Integrated III-V on silicon high-speed transmitter.....	63
III-1. Hybrid III-V on silicon DBR laser	64
III-1.1. Laser overall structure view	64
III-1.2. Adiabatic taper design	66
III-1.3. Reflectors design	70
III-2. Transmitter structure	74
III-2.1. Silicon thickness transition.....	74
III-2.2. Silicon modulator	75
III-2.3. Transmitter layout	76
III-3. Hybrid III-V on silicon transmitter fabrication	77
III-3.1. SOI part fabrication	77
III-3.2. III-V integration.....	79
III-4. Integrated transmitter characterization	84
III-4.1. Stand-alone components.....	84
III-4.2. Complete transmitter characterization.....	89
III-5. Conclusions and improvements	93
Chapter IV: Hybrid III-V on amorphous silicon laser	95
IV-1. Adiabatic coupling into a thin SOI layer	96
IV-1.1. III-V waveguide engineering	96
IV-1.2. Silicon waveguide thickening	98
IV-2. Hybrid III-V on silicon DFB laser design.....	100
IV-2.1. Selected configuration for the DFB laser	100
IV-2.2. DFB laser structure design	101
IV-3. Hybrid III-V on amorphous Si laser fabrication	103
IV-3.1. Amorphous silicon integration schemes	103
IV-3.2. Detailed process flow.....	104
IV-4. III-V on amorphous Si laser characterization	109
IV-5. Conclusions on the use of amorphous silicon.....	111
Chapter V: Conclusions and outlook	113
V-1. Work overview	114
V-1.1. Results synthesis.....	114
V-1.2. Conclusion	115
V-2. Next steps for the transmitter integration	115
V-2.1. Complete fabrication process on 8" or 12" wafers	115

V-2.2. Electro-optical integration	117
V-3. Further improve the components performances.....	117
V-3.1. Laser thermal performances.....	117
V-3.2. Pulse-amplitude modulation (PAM) format.....	119
V-4. Future of silicon photonics	120
List of figures	121
List of tables	127
Bibliography	128
Author list of publications	135
Appendix A: French summary.....	137
A-1. Introduction.....	138
A-2. Modulateur Silicium	144
A-3. Transmetteur hybride III-V sur silicium.....	155
A-4. Laser hybride III-V sur silicium amorphe.....	164
A-5. Conclusions et perspectives	169

Thanks/Remerciements

Après plusieurs années de dur labeur, de sang (en coupant un bout de fromage) et de larmes (pas qu'à cause du bout de fromage), j'ai enfin pu soutenir ma thèse et achever ce manuscrit qui résume au mieux le travail effectué pendant cette période. Malgré tout, cette thèse m'a aussi apporté beaucoup de bonnes choses d'un point de vue personnel, et j'ai vraiment apprécié ces années. Ainsi, j'espère que vous prendrez autant de plaisir à lire ce manuscrit, que j'en ai pris à travailler sur cette thèse (mais peut-être pas à la résumer...). Néanmoins, il serait injuste de dire que tout le mérite m'en revient, car je n'y serais sans doute pas arrivé sans le soutien des nombreuses personnes qui m'ont apporté leur aide tout au long de cette thèse, qui était un véritable travail d'équipe. Je souhaite donc profiter de ces quelques lignes pour les remercier du mieux possible, en espérant n'oublier personne.

Tout d'abord, je souhaite remercier les membres de mon jury, pour l'intérêt porté à mon travail, et d'être venus, bien souvent de loin, pour assister à ma soutenance. Par conséquent, je remercie le professeur Nicolas Grandjean pour avoir accepté la charge de président de jury, ainsi que les professeurs Delphine Marris-Morini et Geert Mortier pour avoir accepté d'exercer la fonction de rapporteurs sur mon manuscrit. Je remercie aussi le professeur Leif Katsuo Oxenløwe et le docteur Di Liang, pour avoir participé au jury de thèse en tant qu'examinateurs. J'ai vraiment apprécié les différents échanges que nous avons eu pendant la soutenance, et j'espère que nous aurons l'occasion de collaborer ensemble dans le futur.

Ensuite, je souhaite remercier mes différents encadrants, qui ont supervisé l'ensemble de mon travail et m'ont apporté leurs précieux conseils tout au long de celui-ci. A commencer par Badhise Ben Bakir, mon encadrant au CEA-LETI, qui m'a initié au monde de la photonique sur silicium, et m'a formé à la fois à la conception, la fabrication et la caractérisation des composants photoniques. Son obstination et sa persévérance m'ont poussé à me dépasser pour atteindre mon but. Je remercie aussi mes encadrants à STMicroelectronics, Alain Chantre et Frédéric Bœuf. Alain est celui qui m'a proposé le premier ce sujet de thèse, et qui m'a accompagné durant la première moitié de celle-ci jusqu'à son départ à la retraite, où Frédéric a ensuite pris le relais. En tant qu'industriels, ils m'ont permis d'obtenir une autre vision de la recherche, que je n'aurais pas pu acquérir si je n'étais resté qu'au CEA-LETI, et ont toujours pris le temps de me conseiller sur la direction à suivre dans mon travail. Pour finir, je remercie aussi mon directeur de thèse, Christian Seassal, qui, même si nous ne sommes pas beaucoup vus à cause de l'éloignement géographique, a toujours été prêt à m'aider en cas de besoin, surtout à la fin de la thèse. C'est grâce à son expérience et à ses judicieux conseils, que j'ai pu achever ce manuscrit et à lui donner sa cohérence, ce qui n'est pas évident sans un œil extérieur.

Je remercie aussi STMicroelectronics, pour m'avoir offert cette opportunité de travailler dans la photonique sur silicium, au sein de l'équipe de Frédéric Bœuf. Je remercie les membres de cette équipe, Charles, Sébastien, Nathalie, Elise, et Aurélie, pour m'avoir offert leurs nombreux conseils avisés lors des réunions hebdomadaires à Crolles, et aussi autour d'un café lors de mes visites occasionnelles. Je remercie aussi Jean-François et Patrick pour les nombreuses discussions que nous avons eu lors de leurs passages au CEA-LETI. Bien entendu je remercie aussi les différents doctorants qui ont débuté avant ou après moi et avec qui il était toujours intéressant d'échanger nos points de vue techniques, ou tout simplement de discuter pour relâcher un peu la pression : Léopold, Boris, Jocelyn (je reviendrais sur vous trois plus tard), Benjamin, Maurin, Jean-Baptiste et Sylvain. Sans oublier Gaspard, qui bien que n'étant pas photonicien, avait débuté sa thèse en même temps que moi, et avec qui j'appréciais beaucoup discuter lors de mes passages à Crolles pour me changer les idées.

Je souhaite aussi remercier le CEA-LETI pour avoir accepté de m'intégrer dans le laboratoire LCPC, et de m'avoir fourni les différents moyens de réussir ce travail de thèse, financé par l'ANR via le programme « IRT Nanoelec ». Je remercie ainsi Sylvie Menezo et Christophe Kopp, tour à tour chefs du LCPC de m'avoir accueilli dans leur équipe de choc, dont les différents membres m'ont toujours apporté leur aide en cas de besoin, et avec qui j'ai passé d'excellents moments. Tout d'abord, je remercie Benjamin, pour m'avoir enseigné les « secrets de la RF », sans lesquels je n'aurais pas pu obtenir d'aussi bons résultats, mais aussi pour les nombreux

échanges que nous avons eu durant toute ma thèse. Je remercie aussi Corrado, qui m'a appris beaucoup de choses sur le monde des III-V qui m'était alors peu connu, et avec qui j'appréciais philosopher sur différents sujets. Je remercie aussi Karim et Marco, qui m'ont toujours incité à persévérer quand je faisais face à des difficultés. Je n'oublie pas non plus Damien, Julie, Sonia, Stéphane (B.) et Paul pour m'avoir appris les ficelles pour gérer mes lots en salle. Bien entendu, je n'aurais pas non plus obtenu mes résultats sans les professionnels de la caractérisation optique du LCPC, que je remercie également : Philippe (qui arrive toujours à trouver un nouvel équipement pour sauver la mise, accompagné d'une bonne blague en-dessous la ceinture), Karen (Gary-the-snail, avec qui j'ai pu partager ma passion pour la culture japonaise), André (qui finira par trouver la poulette), et Olivier (pro de Labview et de la langue française). Merci aussi à Stéphane et Benoît pour nos différentes discussions au niveau du système, qui m'ont permis de voir au-delà des simples composants. Je remercie aussi toutes les personnes du labo avec lesquelles je n'ai pas ou peu travaillé directement, mais avec lesquelles j'appréciais toujours une bonne discussion, scientifique ou non: Vincent, Daivid, Toufiq, Bayram (seul courageux à me suivre au foot !), Thomas, Vincent, Patricia, Laeticia, Laurent, Maryse, Ségolène, Stéphane (M.), Bertrand, Gilles et Jean-Marc. Bien évidemment, que serait le LCPC sans mes camarades doctorants (« on vous exploite, on vous spolie ! ») et avec qui j'ai partagé d'excellents moments ? Je pense bien sûr à ceux dont j'ai partagé le bureau, ainsi que les difficultés liées à nos thèses respectives, et que je remercie pour leur soutien : Léopold, Boris, Jocelyn (oui, les mêmes qu'au paragraphe précédent). Je remercie aussi Antoine et Hélène, avec qui j'ai partagé le même encadrant, ainsi que les mêmes problématiques sur les lasers hybrides. Pour finir, je remercie tous les autres doctorants du laboratoire (Alexis, Olivier, Simon, et Matthieu), ainsi que ceux que j'ai pu rencontrer au sein du CEA-LETI via l'association AITAP (Heimanu, Emilie, Benoît, Lucile, Kévin, Samantha, Pierre, Anna, Elvira, Véronika, et Carine).

Je remercie aussi les différentes membres du DTSI, ingénieurs et techniciens, qui m'ont appris énormément de choses sur les procédés de fabrication en salle blanche et m'ont toujours aidé, lorsque je venais les solliciter pour fabriquer mes composants. Ils sont tellement nombreux que je vais sûrement en oublier, mais je vais quand même essayer : merci donc à Olivier G., Denis, Gilles, Olivier P., Pierre, Ludivine, Philippe, François, Romain, Frédéric, Carmelo, Nicolas, Fabrice, Maurice, Viorel, Sébastien, ... De même, je remercie aussi les membres de l'équipe photonique du LPFE qui ont achevé avec succès la fabrication des lasers et des transmetteurs : Christophe, Tiphaine, et Romain. Bien que je ne puisse pas suivre directement la totalité des procédés de fabrication III-V qui se passaient dans une autre salle blanche, ils ont assuré une fabrication rapide de mes échantillons, tout en me fournissant un excellent suivi de chaque étape avec de nombreuses photos à l'appui.

Bien que je n'y ai passé que peu de temps en comparaison de STMicroelectronics ou du CEA-LETI, je remercie aussi le laboratoire INL et l'Ecole Centrale de Lyon pour m'avoir accepté en tant que doctorant en photonique sur silicium.

A présent, je souhaite clore cette section en remerciant mes proches, à commencer par mes parents, Rossano et Christine, toute ma famille, ainsi que Marion et Sacha, pour avoir cru en moi, et m'avoir soutenu pendant toutes ces années. Si je suis arrivé là où je suis aujourd'hui, c'est bien grâce à vous, et grâce à votre soutien et à votre affection indéfectible.

Pour finir, je te dédie cette thèse Marion, pour m'avoir soutenu tous les jours, même quand je commençais à douter de ma capacité à y arriver. Tu es ce qui m'est arrivé de plus beau pendant cette thèse, et j'ai énormément de chance de t'avoir à mes côtés.

Chapter I: Introduction

The aim of this introductory chapter is to present the context and objectives of this work, focused on silicon photonic integrated transmitters for data transmission in data-centers, as illustrated below. This chapter starts with the general context of this work, as well as a status of photonics nowadays, followed by a more specific presentation of silicon photonic and its current main flaw: the absence of a monolithically integrated optical source. Then, the heterogeneous integration of III-V materials, which is currently the most promising solution for an integrated laser on silicon, will be discussed. Finally, the few existing demonstrations of silicon photonics integrated transmitters are presented, as well as the objectives of this work.



Source: Google.

I-1. Context	2
I-1.1. Information society and photonics	2
I-1.2. Current status of the photonics market	2
I-2. Silicon photonics	3
I-2.1. Why should we use silicon for photonics?	3
I-2.2. Silicon photonics device library	4
I-2.3. Optical source issue in silicon photonics	5
I-3. Heterogeneously integrated III-V on silicon laser	9
I-3.1. General principles of laser operation	9
I-3.2. Hybrid III-V on silicon laser structure description	13
I-3.3. Existing hybrid III-V on silicon laser cavity designs and state-of-the-art.....	17
I-3.4. Advantages of hybrid III-V on silicon lasers and their integration issues	21
I-4. High-speed integrated transmitters	22
I-4.1. State-of-the-art for integrated transmitters in silicon photonics	22
I-4.2. Work objectives and targeted specifications.....	23

I-1. Context

Before entering deeply in the main subject of this work, the general context in which the study has taken place will be presented in this section. Indeed, the world of photonics is vast, and it might be useful to see where to focus our attention.

I-1.1. Information society and photonics

For years, the volume of data exchanged across the world is continuously increasing, and this trend seems unlikely to stop. According to forecasts from Cisco [1], the volume of information transmitted each year via the Internet is now beyond the zettabyte (10^{21}), and, should reach two zettabytes by 2019. This is equivalent to more than 16000 DVDs (4Gbytes) exchanged each second, during one year. With the improvements of Internet access and mobile data rates, even ultra-high-definition (UHD) can be streamed online. Moreover, everything is connected nowadays: smartphones, tablets, watches, televisions, cameras, cars, houses... This multiplication of connected devices (referred as “Internet of Things”) also brings a large share of information to deal with. To manage this amount of information, high data transmission rates are required for distances over several hundreds of meters, especially in the data-centers. Unfortunately, copper-based interconnections are not suited to transmit information over long distances at high frequencies. For instance, a typical coaxial cable will have 1dB/m attenuation at 10GHz, and will lose 90% of its intensity over 10m. When the signal speed increases, the maximum reachable distance decreases, and becomes an issue even for communication inside electrical chips.

To overcome these limitations, alternative ways such as photonics, which make use of photons – rather than electrons – to convey information, are continuously developed. Unlike copper, the attenuation in silica-glass fibers can be as low as 0.2-0.3dB/km. Photonics have already been used for a long time for km-long distance communications, and the minimum distance for its utilization continues to decrease. Another decisive feature of photonics is the wavelength-division multiplexing (WDM), which enables to increase the data transmission rate of a single fiber by using several wavelengths at the same time. However, even if they outperform copper-based interconnections, these devices are quite expensive, thus limiting their use if not absolutely necessary. The large cost of photonic systems can be explained by the materials used in their fabrication (based on InP, GaAs, or LiNbO₃), as well as their packaging, which can become complex especially if discrete components must be assembled in one package.

I-1.2. Current status of the photonics market

The photonic market for communication is vast and can be complex to apprehend. A simple way to classify it is to use a distance-based classification, which is presented here. It can be noted that communications at the scale of the data-centers – which are at the core of this work – are presented with a bit more details.

Very short-reach (below 3 meters)

The term of very-short reach encompasses all connections for distances below one to three meters. It concerns communications between boards, and even inside electronic chips. Currently, this area is not covered by photonics. Indeed, even with their bandwidth limitations, copper-based interconnections can still manage high data rates at those distances, and are still less expensive than photonics. Nevertheless, even if at 40Gb/s copper cable can support signal transmission up to 3m, at 100Gb/s the maximum distance achievable will be 1m. Therefore, the limit distance separating copper and photonics will continue to decrease for higher modulation speeds.

Short-reach (between 3 and 300 meters)

In the short-reach domain (which comprises most of the links in the current data-centers [2]), there is currently one photonic solution which dominates the others: the multi-mode vertical cavity surface emitting laser (MM VCSEL) operating at the 850nm wavelength. This domination over other types of transmitters can be explained by the low-cost fabrication, testing and packaging of the VCSELs. As their name indicates, light being emitted from the surface of the device, they can be easily tested on the wafer before packaging. Their output aperture is also large, which facilitate the coupling with the fiber [3]. Finally, they are also compact and can

benefit from a large-scale fabrication. A single directly modulated VCSEL with a 56Gb/s modulation rate was even recently demonstrated [4], showing that speed is not yet the limitation for these devices. The main issues of the MM VCSEL are the modal and chromatic dispersions, which reduces its maximum range from 300m to 100m , if their data rate goes from 10Gb/s to 25Gb/s [2].

Long-reach (above 300 meters)

Unlike the previous domain, there are no dominating devices in the long-reach area. Transmitters are based on edge-emitting lasers, which can be directly modulated (DML) or externally modulated (EML). Those devices are more expensive than VCSELs since their testing and packaging are more complex, but also because they are generally fabricated on InP wafers, which are smaller than the GaAs wafers used for VCSELs [2]. According to the application, two wavelengths regions are generally used. As it can be seen on **Figure I-1(a)**, the $1.55\mu\text{m}$ region (C-band) corresponds to the minimum of attenuation in silica-glass fibers, and thus it is used for the longest transmission distances. However, additional systems are generally needed to compensate for the dispersion. On the opposite, the $1.3\mu\text{m}$ region (O-band) has a bit more absorption, but has a negligible dispersion, as shown on **Figure I-1(b)**. Therefore, for Datacom applications, which are at most tens of kilometres long, the transmitters generally operate around $1.3\mu\text{m}$, and for telecom applications, which are even longer, they operate around $1.55\mu\text{m}$. In the case of data-centers, the maximum standard reach is often considered to be 10km (with transmitters operating in the O-band), even if several industrials proposed to implement a new standard at 2km [2].

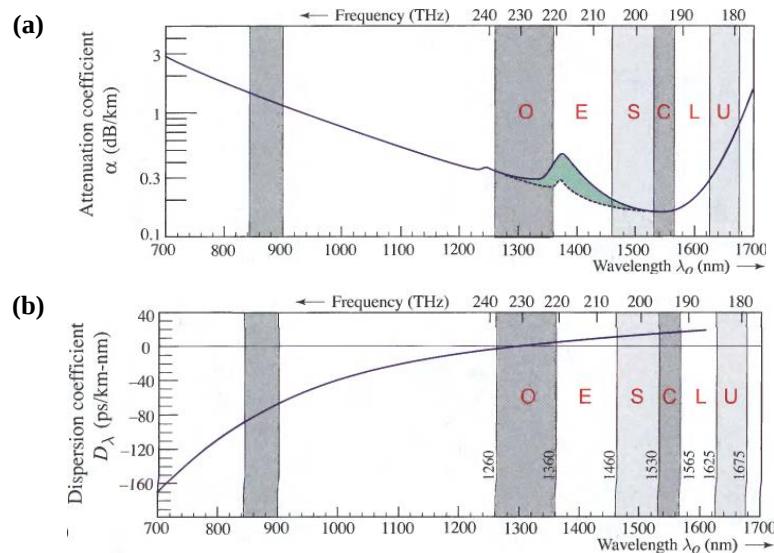


Figure I-1: (a) Attenuation coefficient, and (b) dispersion coefficient dependence on the wavelength for silica-glass fibers. Images from [5].

I-2. Silicon photonics

In this context, silicon photonics is often seen as a challenger which is expected to disrupt the current photonics environment. While it has been in development for more than 30 years, it has known fast developments in the last decade. What can explain this craze for silicon photonics? This section presents the main interests of silicon photonics, its library of components, as well as its main flaw: the lack of an efficient light source.

I-2.1. Why should we use silicon for photonics?

Silicon being transparent above $1.2\mu\text{m}$, it has been investigated for a long time to be waveguide medium for the optical communication wavelengths. The high index contrast between silicon and its native oxide, as well as the buried oxide layer (BOX) present in silicon-on-insulator (SOI) wafers, both providing a tight confinement of the optical mode within silicon waveguides, and enabling efficient and compact optical designs, were another of

its strengths. But silicon photonics attractiveness comes mainly from the prospect of obtaining what was missing for photonics all this time: high performances at low-cost. Indeed, fabrication using the well-controlled silicon processes offers large-scale production capabilities, with high yield and high device reliability. Moreover, silicon photonics does not need the most recent technological developments, needed for transistors with gate length below tens of nm, which should further reduce the cost. Finally, silicon photonics possesses another key advantage, which is the access to integrated drivers and trans-impedance amplifiers (TIA). These circuits can be either monolithically integrated or bonded by flip-chip above the photonic circuit, depending on the fabrication process.

At first, silicon photonics was seen as the way to finally use light for board-to-board or even intra-chip communications. While it is still the case, the first previsions were a bit optimistic, and silicon photonics is currently trying to impose itself in the data-centers, which are its main market. Not only it can have the cost advantage in the long-reach domain, but also in the short-reach, since new mega data-centers are being designed using single-mode fiber (SMF) – instead of multi-mode fiber (MMF) – in order to solve the reach limitation issue of MM VCSELs [2].

I-2.2. Silicon photonics device library

In order to form an optical fiber communication system, several components are needed. An example is displayed on **Figure I-2**, with a system for the 100GBASE-LR4 standard, which is equivalent to four different wavelengths, each modulated at 25Gb/s, for the long reach domain (up to 10km). To emphasize the strength of silicon photonics, the maximum of photonic functions should be transferred in the SOI layer. Thus, in the last decade, silicon photonics has known fast developments to integrate these different components on silicon, leading to a vast device library, with almost all the necessary passive and active components to form silicon photonic circuits.

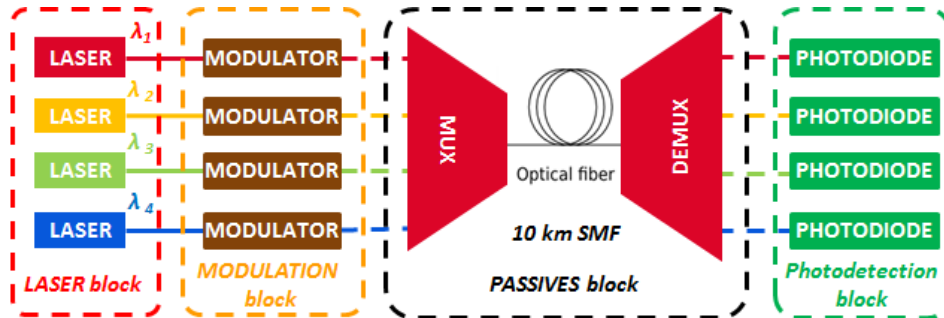


Figure I-2: Example of photonic optical link for the 100GBASE-LR4 standard.

Among the passive components, surface-to-fiber grating couplers have been researched intensively to propose a nearly-vertical coupling with the minimum of optical losses, and simplify the testing and packaging of the devices [6]. Low-loss silicon waveguides are also necessary to realize the circuitry between the different components [7], and the WDM functions are implemented using low cross-talk (de)multiplexers [8]. For the active components, high-speed silicon modulators, used to convert an electrical signal into an optical one, have been demonstrated by several groups (as it will be discussed in **sections II-1.3** and **II-1.4**). As for the photodetectors, since silicon is transparent at the optical communications wavelength, another material is necessary to convert back the signal in the electrical domain. Hopefully, germanium is now a CMOS-compatible material, which can be integrated through epitaxial growth on silicon, in order to form efficient photodiodes [9].

All the components presented in this section have gained a lot in maturity. Now, they start to be fabricated at an industrial level [10], [11], and can also be found in multi-project wafer (MPW) services [12]–[14]. Those silicon photonics platforms differ from one to another, with some proposing to use the same SOI wafer to form the transistors of the driving circuit and the silicon photonics components [10]. However, in recent years, there has been a convergence towards the use of different SOI wafers, either suited for the transistors (with a thin BOX layer below 30nm) or for the silicon photonics components (with a SOI layer thickness generally between 200-300nm, and a BOX layer thickness above 700nm) [11]–[14]. Electronic and photonic circuits can then be connected by wire bonding or three-dimensional flip-chip integration.

I-2.3. Optical source issue in silicon photonics

As explained in **section I-2.2**, almost all the components necessary for optical fiber communication on the silicon photonics platform have already been demonstrated and monolithically co-integrated on the same wafer. However, a last element is still missing: an efficient laser source. As it will be explained in the **section I-3.1**, a laser needs an amplifying medium to operate. However, due to its band structure, silicon is a poor candidate to bring both photon emission and amplification, necessary to form the laser. The reasons why are explained in this section. The solutions envisaged to monolithically integrate a laser in silicon are also presented, as well as the solutions used nowadays to couple an external optical source to the silicon photonic circuits.

Light amplification in semiconductors and influence of the band structure

In the case of semiconductors, amongst the several electronic transitions resulting from the interaction between electrons and photons, three interband transitions are generally considered for photonics devices and schematized on **Figure I-3**. The first one is the spontaneous emission of a photon, due to the recombination of an electron in the conduction band with a hole in the valence band. For this phenomenon to happen, energy must be conserved during the recombination, and the photon energy ($h\nu$) must be larger than the bandgap energy (E_g). Thus, for an electron occupying an energy level E_2 , and a hole occupying an energy level E_1 , the following condition must be fulfilled:

$$E_2 - E_1 = h\nu > E_g \quad (\text{I.1})$$

Where h is Planck's constant and ν is the photon frequency. The emission time and direction being random, the radiation field is not coherent. It is the main mechanism used in light-emitting diodes (LED). The second one is the absorption of a photon, where the photon energy generates an electron in the conduction band, while leaving a hole in the valence band. Photodiodes are based on this mechanism.

The last one is the stimulated emission of a photon, where an incident photon stimulates the recombination of an electron in the conduction band with a hole in the valence band, while generating a new photon. Semiconductor optical amplifiers (SOAs) and lasers are based on this amplification mechanism (laser being the acronym for Light Amplification by Stimulated Emission of Radiation), where the new photon is identical to the incident one, resulting in coherent emission.

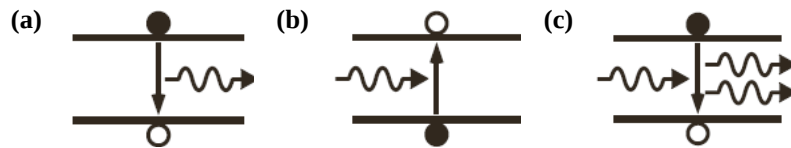


Figure I-3: Interband transitions between the conduction (top) and valence (bottom) bands, in which the energy is supplied by or given to a photon: (a) spontaneous emission, (b) absorption, (c) stimulated emission. Filled circles represent the electrons, and open circles the holes. Images from [3].

Stimulated emission and absorption are competing effects, both depending on the density of electrons-holes pairs available for recombination. Since electrons are naturally present in the valence band, a “pump” is required to populate the conduction band, in order to overcome the absorption phenomenon and reach optical gain. This state is reached when the carrier population has been inverted, and the following condition is satisfied:

$$E_{Fc} - E_{Fv} > h\nu > E_g \quad (\text{I.2})$$

Where E_{Fc} and E_{Fv} are respectively the conduction bands and valence bands quasi-Fermi levels. An optical pumping can be used to generate them through absorption, but semiconductor lasers present the advantage to have access to an efficient electrical pumping scheme via current injection in a p - n junction. The structures are then referred as laser diodes. This solution permits to increase the system integration, since an additional optical source is not needed, but these lasers are also more complex to operate compared to optically pumped ones. In both optical and electrical pumping, two types of operations are considered: pulsed and continuous. Pulsed operation is used to avoid self-heating effects, which are detrimental for the laser performances as explained later

in **section I-3.1**. Continuous operation is more difficult to reach, but also simplifies greatly the laser driving scheme.

However, spontaneous and stimulated emission mechanisms are unlikely to happen in indirect band structures such as silicon. Indeed, those electronic transitions require the conservation of both energy and momentum. Since photons carry a large energy, but a negligible momentum, the momenta of electrons and holes used during the radiative recombination must be approximately equal. This condition is fulfilled relatively easily in the case of a direct band structure material such as InP, where the lowest energy levels of the conduction and valence bands are aligned as shown on **Figure I-4**. However, it is not the case in an indirect band structure material such as silicon, where electrons in the conduction band are not aligned with holes in the valence band. Therefore, phonons (which carry a large momentum, but a negligible energy) must be involved to satisfy the momentum conservation condition. Since a phonon-assisted transition requires at the same time electrons, photons and phonons, its probability of occurrence is extremely low. Intraband transitions (such as free-carrier absorption) or non-radiative electron-hole recombination (such as Auger recombination) are more likely to happen. Therefore silicon is not suited to form an optical source, and alternative solutions must be used to couple light in the silicon photonic circuit.

It can be noted that even though interband absorption also requires the conservation of both energy and momentum, it is not unlikely in indirect bandgap structures. Indeed, the electron can transfer its momentum to a phonon after being generated in the conduction band, thus the three bodies (electrons, photons and phonons) are not needed at the same time. The same phenomenon happens to the generated hole in the valence band.

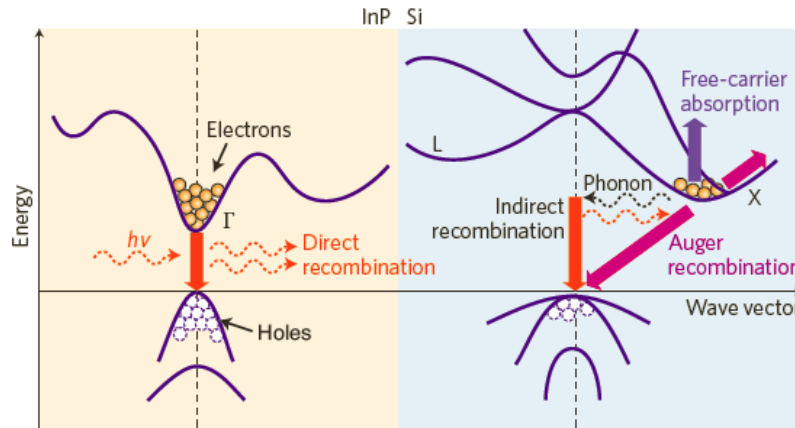


Figure I-4: Energy band diagrams and major carrier transition processes in a direct band structure such as InP (left) and in an indirect band structure such as silicon (right). Image from [15].

Monolithically integrated laser sources on silicon

Even though silicon is not suited to form efficient light sources, research still goes on to form full-silicon lasers. Thus, in 2005, an optically pumped in continuous operation silicon laser based on Raman scattering was reported [16]. However, even if amplification is possible using Raman scattering, this effect requires external signal and pump beams to operate [15], [17], which reduces the system integration level.

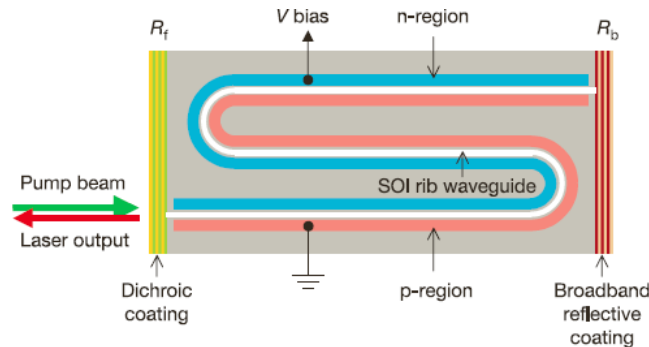


Figure I-5: Schematic view of a Raman scattering laser. Image from [16].

Another considered option is to use Ge as a gain material. As explained in **section I-1.2**, Ge is already compatible with CMOS fabrication processing and used to form photodiodes in silicon photonic circuits. It has an indirect band structure, but unlike silicon (depicted on **Figure I-4**), the energy gap between the top of the valence band and the same momentum valley in the conduction band (0.8eV) is close to its actual indirect bandgap (0.66eV), as it can be seen on **Figure I-6(a)**. A direct bandgap of 0.8eV is equivalent to a bandgap wavelength of $1.55\mu\text{m}$, which makes it interesting for telecommunication applications. By applying tensile strain to the Ge, both direct and indirect bandgaps are reduced, and the effect is even stronger for the direct bandgap [18], as shown on **Figure I-6(b)**. Strain can be induced during the growth on silicon, due to the difference of thermal expansion coefficient between Si and Ge. However, strain will also increase the bandgap wavelength above the standards for telecommunication. A proposed compromise was to use limited tensile strain and free electrons brought by heavy n-doping to fill the indirect bandgap valley, in order to increase the probability of occupation of the direct bandgap valley by injected carriers [19], as depicted on **Figure I-6(c)**. Using these methods, an electrically pumped Ge-on-Si laser in pulsed regime has been announced in 2012 [20]. However, those results have been reproduced only once since then [21], and there are still contradictory opinions on this structure having reached or not laser operation [22]. Therefore, a lot of work is still ongoing to obtain a robust Ge-based laser. Current approaches are based on using further increasing the tensile strain (by using micromechanical structures [22] or external stressor [23]), or GeSn alloys (which are equivalent to apply a tensile strain) are promising, but the bandgap wavelength would be over $2\mu\text{m}$ and thus would be more dedicated to mid-infrared laser applications such as sensing rather than telecommunications.

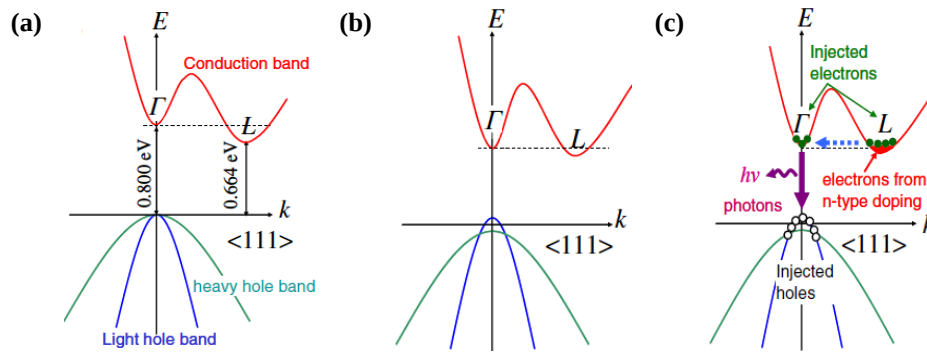


Figure I-6: Band structures of (a) bulk Ge, (b) tensile strained intrinsic Ge, and (c) tensile strained n-doped Ge. Images from [18].

Finally, another option for a monolithically integrated laser source on silicon is the epitaxial growth of materials with a direct band structure, such as GaAs and InP, which are commonly used to form laser diodes. The growth of III-V materials has recently known a renewed interest of the microelectronics industry, for the possibility to use them to form high-mobility compound semiconductors in the next-generation of CMOS transistors [24]. However, epitaxial growth of GaAs and InP on silicon is not a trivial task, due to a large lattice mismatch of respectively $\approx 4\%$ and $\approx 8\%$, as well as large differences on their thermal expansion coefficients of respectively $\approx 120\%$ and $\approx 80\%$ [15]. Such mismatches result in a large density of defects, such as misfit or threading dislocations, which would be detrimental to the electro-optic performances of the devices. The most common approaches are to use micrometres-thick buffer layers and/or strained superlattices to accommodate the mismatches, before growing the active layers for the laser. Based on these techniques, an electrically pumped InAs/GaAs grown on silicon laser in continuous operation and emitting in the $1.3\mu\text{m}$ wavelength region was recently demonstrated, and displayed on **Figure I-7(a)** [25]. However these thick buffer layers between the silicon layers and the active layers complicate the light coupling between the laser source and the silicon photonic circuit. Due to their topography, it will be also difficult to integrate them in a complete fabrication process with the other silicon photonic devices. Nevertheless, epitaxial growth is still improving, and a high-quality InP layer has been recently grown directly on silicon, with dislocations limited to a 20-nm -thick layer. An optically pumped laser under pulsed operation based on this directly grown InP layer has also been demonstrated, emitting around the InP bandgap wavelength ($\approx 0.92\mu\text{m}$), as shown on **Figure I-7(b)** [26]. Those results are really promising, but a lot of work is still needed to obtain electrically pumped lasers, shift the emission wavelength in the telecommunication range, and couple it in the silicon photonic circuit.

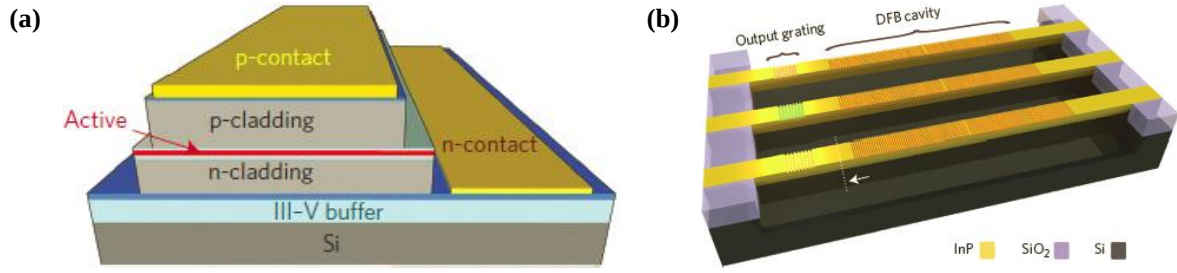


Figure I-7: Schematic view of InP directly grown on silicon laser (a) with a thick III-V buffer [25] and (b) without buffer [26].

Externally integrated laser sources

Though a lot of progresses have been realized towards monolithically integrated laser on silicon in the last few years, the current solutions are still not mature enough to be integrated with the silicon photonic circuits. Therefore, the current transmitters in silicon photonics generally rely on externally integrated laser sources. Those sources are generally pre-processed III-V laser diodes, which can be integrated on top of the silicon photonic circuit by using techniques such as flip-chip bonding [27], [28]. Examples of such externally integrated laser source are displayed below, where the source can be integrated in a hermetic package, as on **Figure I-8(a)**, or as a bare chip, as on **Figure I-8(b)**.

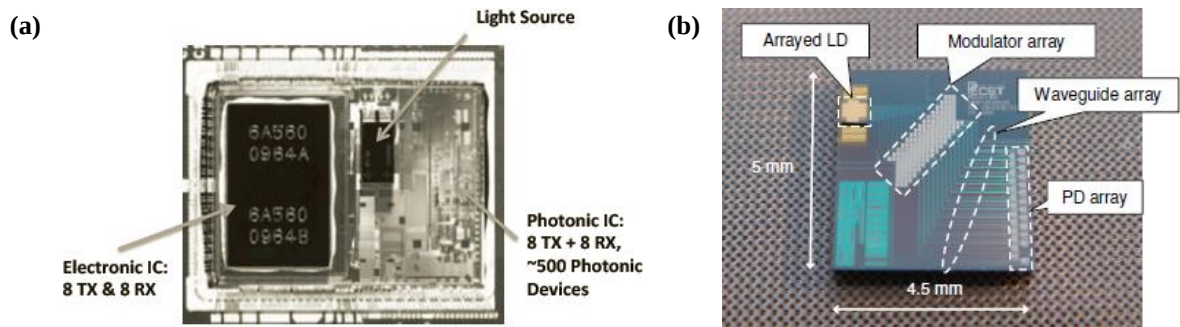


Figure I-8: Examples of silicon photonic circuits with external laser diodes (referred as LD) integrated to the silicon substrate. (a) LD in a hermetic micro-package [27], (b) Arrayed LD chip [28].

Different methods exist in order to couple light from the external laser diodes into the silicon photonic circuit. Thus, systems based on additional ball lenses and prisms [29], or lenses and mirrors integrated in the laser diode [30] have been proposed to offer a near-vertical coupling between the edge-emitting laser diode and a fiber-to-surface grating coupler in the SOI layer, as depicted on **Figure I-9(a)-(b)**. Another possibility is to use an edge-coupling scheme between the edges of the laser diode and the silicon waveguide. In order to accommodate the modal shape difference between the output of the laser diode and the input of the silicon waveguide – thus reducing the coupling losses –, three common schemes can be found in the literature. The first one is to use a thin SOI (≈ 200 to 300-nm -thick) fabrication platform, with specifically designed silicon taper waveguides cladded in a thick SiO_2 layer [31]–[35], as displayed on **Figure I-9(c)**. The second one is based on a thicker SOI ($\approx 3\text{-}\mu\text{m}$ -thick) fabrication platform [36], where the other silicon photonic components are also defined [37], as shown on **Figure I-9(d)**. The last one makes use of a thick low-index layer such as SiON [38], [39], or SiN_x [40], [41], cladded in SiO_2 to form an edge-coupler, and then to couple the light in a thin SOI waveguide, as depicted on **Figure I-9(e)**.

While an externally integrated laser offers a solution to the absence of a monolithically integrated light source in silicon photonic circuits, it also presents several inconveniences. The proposed coupling schemes bring non-negligible coupling optical losses, between -1.6dB and -4dB for the state-of-the-art [29], [31]–[36], [38]–[41]. These coupling losses result in a need to increase the power from the laser, thus in the rise of the overall power consumption. Additionally, they also have a limited tolerance toward spatial misalignments. The surface misalignment limitation for -1dB additional coupling losses is generally below $\pm 2\mu\text{m}$ [29], [32]–[35], [38]

which requires accurate flip-chip bonding techniques. These techniques are generally time-consuming, and reduce the throughput for a large-scale production. Thus, it will also increase the cost and complexity of the complete silicon photonic circuit packaging.

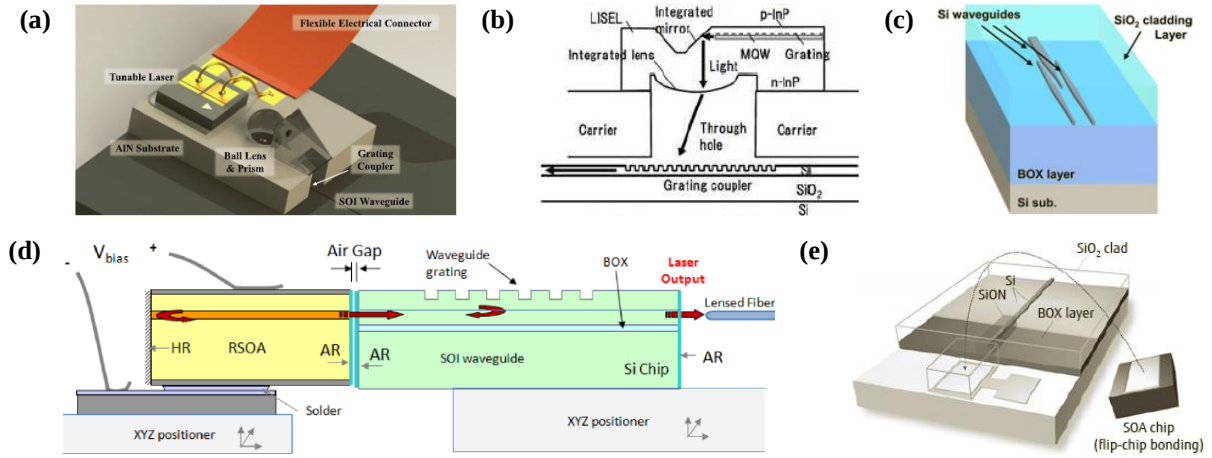


Figure I-9: Coupling schemes between an edge-emitting laser diode and a silicon photonic circuit. Near vertical coupling using (a) additional ball lens and prism [29], and (b) integrated mirror and lens [30]. Edge-coupling using (c) specifically shaped thin SOI waveguides [32], (d) thick SOI waveguides [36], and (e) thick SiON edge-coupler and a thin SOI waveguide [39].

Therefore, the development of a competitive laser source integrated at the wafer-scale, which would not suffer from these issues, is a major objective for silicon photonics. Since monolithically integrated lasers are not yet available, other options must be considered. Currently, the most promising solution for a wafer-level laser source is the heterogeneous integration of III-V materials via bonding, which is described in the next section.

I-3. Heterogeneously integrated III-V on silicon laser

The aim of this section is to introduce the principles of heterogeneously integrated laser sources. These sources rely on III-V stacks formed by epitaxial growth on their native substrate (InP), and integrated on silicon at the wafer-level to obtain the gain necessary to form the laser. These types of lasers are intrinsically different from the externally integrated lasers presented in **section I-2.3**, even with those which feedback is provided by components in the SOI layer [35], [36], [38]–[41]. Indeed, in those cases, the gain material is pre-processed with etching and metallization steps (thus ready for direct use), while it is not the case for the lasers presented in this section, where those steps are realized after bonding, directly on the SOI layer. Therefore, they do not suffer from the same drawbacks, such as a low tolerance towards spatial misalignments.

Before describing the heterogeneously integrated III-V on silicon laser, this section will start with a quick reminder on the general principles of laser operation and on the specific figure of merits (FOM) that will be needed in this document to evaluate the lasers performances. After this introduction, the structure and main operation principles of the hybrid device are detailed. Then, a state-of-the-art of the hybrid III-V on silicon laser is presented. Finally, the advantages of such laser sources are listed, as well as the issues that must be solved to complete their integration in silicon photonics fabrication platforms.

I-3.1. General principles of laser operation

Basically, a laser diode is a SOA converted in an optical oscillator by using a path with optical feedback, as shown on **Figure I-10**. An input light signal (provided by spontaneous emission in the case of laser diodes) is amplified, and the output returns to the input through the feed-back loop, to be amplified again. The process continues until the amplifier gain saturates, and a steady-state is reached.

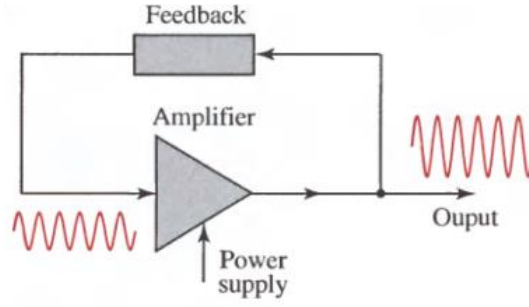


Figure I-10: Schematic view of an optical oscillator. Image from [5].

The general principles of a laser diode can be explained as the following. We will consider the case of light propagating in a waveguide fundamental transverse electric (TE) mode, along the z -axis. Its electric field can be written as:

$$\vec{E}(x, y, z, t) = \vec{e} E_0 U(x, y) e^{j(\omega_0 t - \tilde{\beta} z)} \quad (\text{I.3})$$

Where $\vec{e}$ is the unit vector indicating the TE polarization, E_0 is the magnitude of the field, $U(x, y)$ is the normalized transverse electric field, ω_0 is the optical wave angular frequency, t is the time, and $\tilde{\beta}$ is the complex propagation constant. This constant includes any loss or gain phenomenon, thus is defined as:

$$\tilde{\beta}(\lambda) = \frac{2\pi\tilde{n}(\lambda)}{\lambda} + \frac{j}{2} [\tilde{g}(\lambda) - \tilde{\alpha}(\lambda)] \quad (\text{I.4})$$

Where λ is the optical wavelength, $\tilde{n}$ is the mode effective index of refraction, $\tilde{g}$ is the transverse modal gain, and $\tilde{\alpha}$ is the transverse modal loss. The transverse modal gain parameter refers to the competition between the mechanisms of absorption and stimulated emission described in **section I-2.3**, and can be either positive or negative, depending if enough free carriers have been injected to reach population inversion. If the gain coefficient (g) is constant in a restricted region, and zero elsewhere, the transverse modal gain can be noted as $\tilde{g} = \Gamma_{xy} g$, where Γ_{xy} is the transverse confinement factor in the gain region. The transverse modal loss coefficient takes into account all mechanisms such as free-carrier absorption, and scattering due to the waveguide roughness.

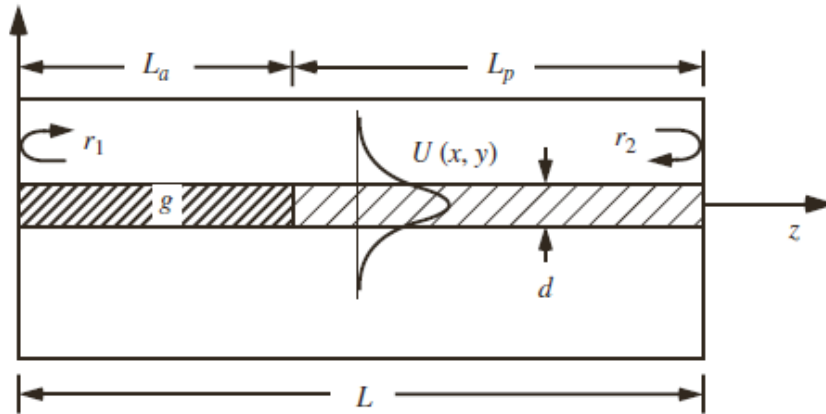


Figure I-11: Schematic view of a laser cavity, with an active region (with a length L_a) and a passive region (with a length L_p). The gain region is restricted to the dashed area of the active region. Feedback is provided by two mirrors with respective complex reflection coefficients r_1 and r_2 . Image from [3].

Next, we will consider the case of a typical laser cavity of length L , with an active region of length L_a , and a passive region of length L_p , as depicted on **Figure I-11**. The passive region is transparent by definition, so in this region $g = 0$. Thus, the propagation constants of the active and passive region can be respectively noted as $\tilde{\beta}_a = \frac{2\pi\tilde{n}_a}{\lambda} + \frac{j}{2} (\tilde{g}_a - \tilde{\alpha}_a)$ and $\tilde{\beta}_p = \frac{2\pi\tilde{n}_p}{\lambda} - \frac{j}{2} \tilde{\alpha}_p$. The necessary feedback is provided by two mirrors with

respective complex reflection coefficients $r_{1,2}(\lambda) = \rho_{1,2}(\lambda)e^{j\varphi_{1,2}(\lambda)}$. At least one of the mirrors must have a partial reflectivity, so that the amount of transmitted light will be coupled into the photonic circuit. Therefore, the ratio between the electric fields before and after one round trip in the cavity can be written as:

$$\frac{E(x,y,2L,t)}{E(x,y,0,t)} = r_1 r_2 e^{-2j\tilde{\beta}_a L_a} e^{-2j\tilde{\beta}_p L_p} \quad (\text{I.5})$$

The laser will start to oscillate when the ratio of electric field will be equal to 1. This state called laser threshold gives two conditions on the amplitude **Eq. (I.6)** and the phase **Eq. (I.7)**, which are:

$$\rho_1 \rho_2 e^{(\Gamma_{xy} g_{th} - \tilde{\alpha}_a) L_a} e^{-\tilde{\alpha}_p L_p} = 1 \quad (\text{I.6})$$

$$\varphi_1 + \varphi_2 - \frac{4\pi\tilde{n}_a L_a}{\lambda} - \frac{4\pi\tilde{n}_p L_p}{\lambda} = 2m\pi \quad (\text{I.7})$$

Where the “th” subscript refers to the threshold value and m is an integer. For clarity, **Eq. (I.6)** can be written as:

$$\Gamma g_{th} = \frac{\tilde{\alpha}_a L_a + \tilde{\alpha}_p L_p}{L} + \frac{1}{L} \ln\left(\frac{1}{\rho_1 \rho_2}\right) \quad (\text{I.8})$$

Where $\Gamma = \Gamma_{xy} \frac{L_a}{L}$ is the gain region confinement factor. Using **Eq. (I.8)**, it can be seen that the gain threshold value corresponds to the total sum of the propagating losses in the cavity and the mirror losses. Thus, if enough free carriers are injected so that the modal gain overcomes the total loss at each round-trip the amplitude condition for lasing will be fulfilled. The gain coefficient being wavelength-dependent, laser oscillation can be reached for several wavelengths. However, it will only occur at the few ones which also respect the phase condition in **Eq. (I.7)**. These lasing wavelengths are also referred to as cavity modes. To push further the selection and reach single wavelength operation, a mode selection filter can be added to the cavity, as depicted on **Figure I-12(a)**. For instance, if the entire gain spectrum shown in **Figure I-12(b)** is above the threshold value, instead of having all the cavity modes fulfilling the laser oscillation condition, only the one chosen by the mode selection filter becomes a lasing mode.

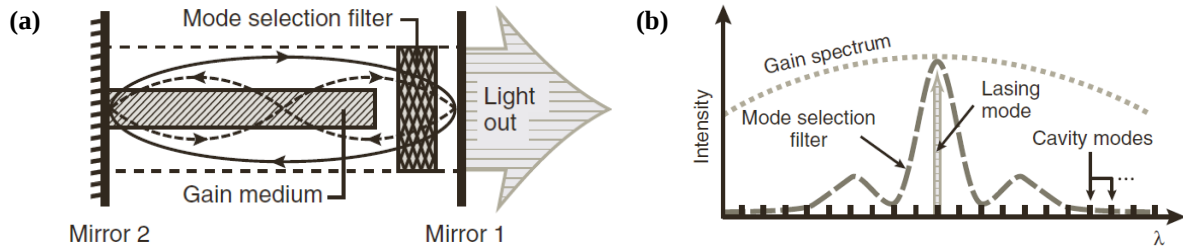


Figure I-12: (a) Schematic view of the different elements needed to form a laser for single wavelength operation. (b) Spectral characteristics of each element. Images from [3].

The spacing between the cavity modes is referred as the free spectral range (FSR). If the mirrors phase is equal to zero (for instance in the case of cleaved facets as mirrors), the FSR can be expressed as [3]:

$$FSR = \frac{\lambda^2}{2(\tilde{n}_{ga} L_a + \tilde{n}_{gp} L_p)} \quad (\text{I.9})$$

Thus, the FSR will be smaller for longer cavities, and more cavity modes will satisfy the phase condition for lasing. **Eq. (I.9)** is not accurate for cavities with more complex mirrors (such as gratings), but still gives a good approximation for the wavelength spacing.

Figures of merit

Several figures of merit (FOM) can be used to evaluate the performances of the laser. Amongst them, only a few (all static) are presented in this document. Dynamic characteristics as the frequency response and large-signal response have not been investigated for the lasers, but are discussed in the next chapter. These FOM can be extracted from three typical plots: the optical output power (P) versus the driving current (I_d), the voltage (V_d) versus the driving current, and the output spectrum.

An example of a typical $P(I_d)$ plot is shown on **Figure I-13(a)**. When the current injected in the junction increases, only few photons are generated at first by spontaneous emission, and the collected optical power is small. When the threshold condition is reached, the laser oscillation starts and the optical output power increases drastically. The corresponding value of current is the current threshold (I_{th}), which should be as low as possible to reduce the electrical consumption of the laser. Since it also depends on the gain section geometry, it is also useful to provide the current threshold density (J_{th}), which is given by the ratio of I_{th} and the surface of the active gain region. The maximum power (P_{max}) can also be extracted this way. The slope $\Delta P / \Delta I_d$ of the $P(I_d)$ plot above threshold is also helpful to determine the energetic efficiency of the laser. It can be used as it is, or as differential quantum efficiency (η_d):

$$\eta_d = \frac{q}{h\nu} \frac{\Delta P}{\Delta I_d} \quad (\text{for } I_d > I_{th}) \quad (\text{I.10})$$

It can be seen that above a certain current value, the output optical power starts to decrease. This roll-off is attributed to self-heating effects (partially due to Joule heating), which increases with the injected current. The temperature increase is detrimental for the lasers, since it will enhance mechanisms such as Auger recombination and carrier leakage outside the active region before stimulated recombination, thus reducing the conversion efficiency. In order to limit the Joule heating, a straightforward way is to have a low series resistance (R_s). This resistance can be extracted from the above threshold slope of the $V(I_d)$ plot displayed on **Figure I-13(b)**. The $V(I_d)$ plot gives also the electrical power (P_{elec}) used to drive the laser, which can be used to extract the wall-plug efficiency (WPE):

$$WPE(I_d) = \frac{P(I_d)}{P_{elec}} = \frac{P(I_d)}{V_d I_d} \quad (\text{I.11})$$

Finally, the laser spectral characteristics are also studied. An example of output spectrum plot at a fixed driving current is shown on **Figure I-13(c)**. The peak wavelength (λ_p) is defined as the wavelength with the highest emitted optical power over the spectrum. The side-mode suppression ratio ($SMSR$) is used to determine if single-wavelength operation is reached or not. It is defined as:

$$SMSR \text{ (dB)} = 10 \log_{10} \left(\frac{P(\lambda_p)}{P(\lambda_1)} \right) \quad (\text{I.12})$$

Where λ_1 is the wavelength with the second-highest emitted optical power over the spectrum. Single-wavelength regime is generally considered as reached for a $SMSR$ above 30dB.

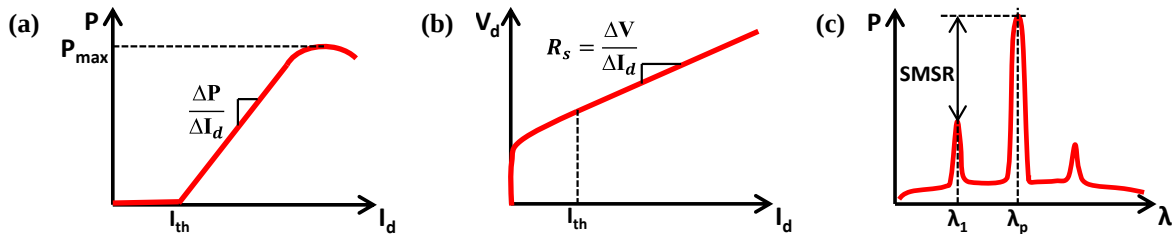


Figure I-13: Typical plots used to characterize the laser diode. (a) Optical output power and (b) voltage versus drive current. (c) Laser output spectrum at a fixed driving current.

Now that the principles of laser operation have been reminded, and the useful metrics have been defined, we will introduce a wafer-level heterogeneously integrated laser source on silicon.

I-3.2. Hybrid III-V on silicon laser structure description

Heterogeneous integration

As explained in **section I-2.3**, there are currently no short-term solutions for a monolithically integrated laser source on silicon. However, an alternative way to obtain the optical gain necessary for a laser is to heterogeneously integrate a gain region (usually made of III-V materials) via bonding. Three types of bonding have been demonstrated in the literature up to now: direct (or molecular) bonding [42]–[58], adhesive bonding [59]–[62], and metallic bonding [63]–[68]. The direct bonding technique relies on a thin oxide layer between the SOI and III-V wafers, which thickness can vary from a few nanometers to a hundred nanometers. The surface of the oxide layer must be as clean and smooth as possible to ensure a perfect bonding without defects, and may require a chemical-mechanical polishing (CMP) step. On the other hand, adhesive bonding uses a polymer layer such as Benzocyclobutene (BCB) between the SOI and III-V wafers, which relaxes the requirements toward the surface roughness. However, such polymer layer is not commonly used in CMOS fabrication platform, and also have a low thermal conductivity, which may prevent the laser from efficient heat dissipation [69]. Contrarily to the other two bonding process, metallic bonding offers at the same time the advantages of relaxed requirements on the bonding surface, but also an excellent thermal conductivity. However, this metallic layer can cause a large optical absorption if too close to the optical mode. To solve this issue, methods such as selective area metal bonding can be used [63]–[65], but require an accurate alignment during the bonding step. Therefore, the direct bonding is the most commonly used to heterogeneously integrate III-V materials on silicon.

The gain region can be bonded on the SOI wafers as entire wafers [70] or as individual dies [70]–[73]. An example of III-V wafer bonded on SOI is displayed on **Figure I-14(a)**. While wafer bonding presents the advantage of integrating III-V materials over a large surface in a single bonding step (with a good yield due to the improvements of the bonding techniques), the size of the III-V wafers are currently limited to 6" at most (but commercially available wafers are generally limited to 3") and cannot cover the entire surface of the SOI wafers (8" or 12"). Those wafers are also expensive, but the majority of the bonded material will be etched away during the processing. These issues can be solved by the bonding of individual III-V dies on the silicon wafer, as shown on **Figure I-14(b)**. This way, a single III-V wafer can cover the complete SOI wafer. However, the bonding yields, as well as the throughput, still need to be improved.

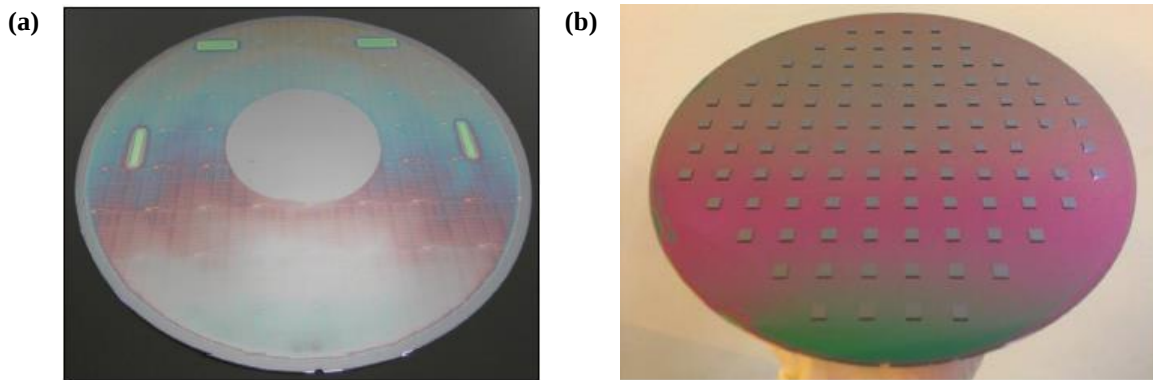


Figure I-14: (a) 3" III-V wafer bonded on an 8" SOI wafer. (b) 104 III-V dies (with dimensions of $5\text{mm} \times 5\text{mm}$) bonded on an 8" SOI wafer. Second image taken from [70].

Once the III-V materials have been bonded, they are processed (etching and metallization steps) on the SOI wafer to form a hybrid III-V over silicon laser. The composition of the bonded III-V layers can vary, but generally have a thick ($\approx 1\text{-}2\mu\text{m}$) *p*-doped InP layer, an active gain region, and a *n*-doped InP layer. Details on the composition are presented later in **section III-1.1**. After bonding, a waveguide is formed by etching, and metallic contact layers are deposited to drive electrically the laser diode. An example of a typical hybrid structure, constituted of a III-V mesa and a silicon waveguide, is presented on **Figure I-15**.

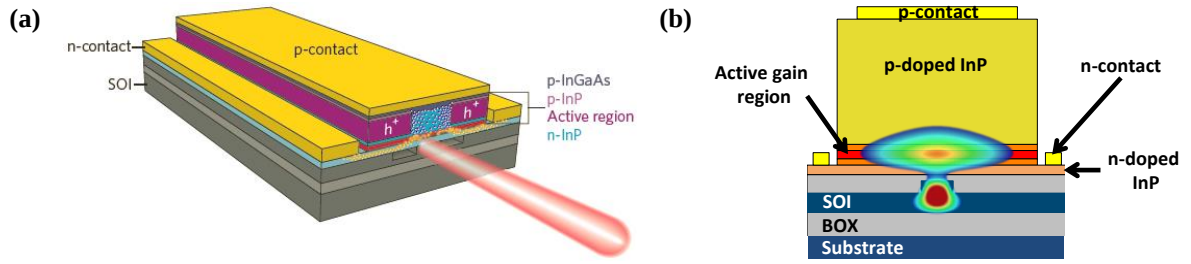


Figure I-15: Schematic views of a hybrid III-V over silicon structure. (a) Three-dimensional view [15], and (b) Cross-sectional view, with the optical mode distributed between the active and passive waveguides.

Within this hybrid structure, we can retrieve the different elements necessary to form laser diode with single-wavelength operation presented on **Figure I-12(a)**. The active region brought by bonding will provide the necessary gain if electrically pumped. The feedback mechanism as well as the mode selection filter can both be realized in the III-V or silicon regions. The most common schemes are detailed later in **section I-3.3**. For the output coupling, as it can be seen on **Figure I-15(b)**, the optical power is shared between the active and passive waveguides, and the percentage of power in each of them depends on their geometry. In some hybrid lasers designs, the optical mode is mostly confined in the silicon waveguide, while the rest is evanescently coupled into the active region [42]–[45], [47]–[49], [55], [56], [61]. This way, light can be easily coupled into the silicon photonic circuit, but at the cost of a lower modal gain. A thin bonding layer is also required to have enough gain. The opposite design, with a large confinement in the active region and an evanescent coupling in the silicon waveguide has also been proposed [54], [74]–[77], but essentially for micro-disks structures which will be presented later in **section I-3.3**.

The best solution to solve this trade-off is to have the light confined in the active region at first, to reach a large modal gain, and then to transfer it into the silicon waveguide at the edges of the hybrid structure, to easily couple it in the silicon photonic circuit. To realize this power transfer, we will rely on the coupled-mode theory.

Applications of coupled-mode theory to the hybrid waveguide structure

First, we need to consider two isolated waveguides (1 and 2). For simplification, we will only consider their fundamental TE modes, with respective effective refractive index $\tilde{n}_1$ and $\tilde{n}_2$. We will also consider that the waveguides are lossless, meaning that their propagation constants become $\tilde{\beta}_{1,2} = 2\pi\tilde{n}_{1,2}/\lambda$. Therefore, using **Eq. (I.3)** the total electric field can be expressed as:

$$E(x, y, z) = E_1(z)U_1(x, y)e^{-j\tilde{\beta}_1 z} + E_2(z)U_2(x, y)e^{-j\tilde{\beta}_2 z} \quad (\text{I.13})$$

In the case of two completely isolated waveguides, their field magnitudes remain constants along the z -axis. However, if they are close enough so that their electric field distribution $U_{1,2}$ overlaps each other, each waveguide will act as a perturbation for the other, and their field magnitudes will vary along the z -axis. In the case of weak coupling, several approximations can be made, leading to two coupled differential equations referred as the coupled-mode equations, which can be found in [3], [5]. To solve these equations, one can use the supermodes formalism, which has been applied to hybrid waveguides structures in [78]–[80]. To do so, the total electric field is presented as a column vector in the basis of $U_1(x, y)$ and $U_2(x, y)$ as:

$$\hat{E}(x, y, z) = \begin{bmatrix} E_1(z)e^{-j\tilde{\beta}_1 z} \\ E_2(z)e^{-j\tilde{\beta}_2 z} \end{bmatrix} \quad (\text{I.14})$$

Two solutions for the total electrical field can be found from the coupled-mode equations. These two solutions are referred as even (E_e) and odd (E_o) supermodes, and are expressed as:

$$\hat{E}_e(x, y, z) = \frac{1}{\sqrt{2}} \begin{bmatrix} \sqrt{1 - \delta/S} \\ \sqrt{1 + \delta/S} \end{bmatrix} e^{-j(\bar{\beta}+S)z} = \hat{e}_e e^{-j\tilde{\beta}_e z} \quad (\text{I.15})$$

$$\hat{E}_o(x, y, z) = \frac{1}{\sqrt{2}} \begin{bmatrix} -\sqrt{1 + \delta/S} \\ \sqrt{1 - \delta/S} \end{bmatrix} e^{-j(\bar{\beta}-S)z} = \hat{e}_o e^{-j\tilde{\beta}_o z} \quad (\text{I.16})$$

With:

$$\delta = \frac{\tilde{\beta}_2 - \tilde{\beta}_1}{2} \quad (\text{I.17})$$

$$\bar{\beta} = \frac{\tilde{\beta}_2 + \tilde{\beta}_1}{2} \quad (\text{I.18})$$

$$S = \sqrt{\delta^2 + \kappa^2} = \frac{\tilde{\beta}_e - \tilde{\beta}_o}{2} \quad (\text{I.19})$$

Where $\tilde{\beta}_e$ and $\tilde{\beta}_o$ are the respective propagation constants of the even and odd supermodes, and κ is the coupling strength between the two waveguides, which increases with the overlap of U_1 and U_2 . From **Eqs. (I.15)-(I.16)**, we can deduce three cases that are summarized in **Figure I-16**. From these three cases, two types of coupling can be realized between the waveguides: directional and adiabatic.

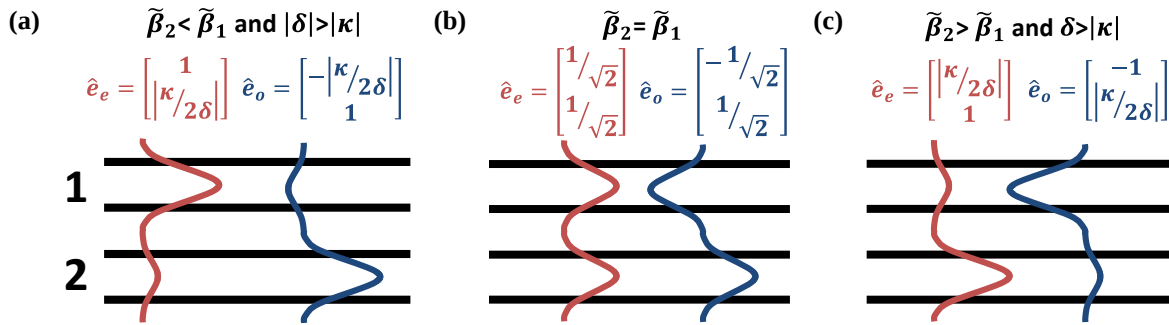


Figure I-16: Limit cases of waveguide coupling where: (a) the even/odd supermode is mostly confined in waveguide 1/2, (b) both supermodes are equally shared between the two waveguides (phase matching), and (c) the even/odd supermode is mostly confined in waveguide 2/1.

The first one is based on the case described in **Figure I-16(b)**, which is called the phase matching condition. In this case, each supermode is equally distributed in each waveguide, and their field magnitudes do not vary along the z -axis. However, their respective propagation constants are different. Thus, when propagating in the waveguides, they will periodically have identical or opposite phases, and alternatively interfere constructively or destructively in each separated waveguide. For instance, if they have the same phase at $z = 0$, all the power will be located in waveguide 2, and after propagating for a distance $L_{dc} = \pi/2\kappa$, their phase difference will be equal to π – according to **Eq. (I.19)** –, and the power will be completely transferred to waveguide 1. While directional coupling is an excellent way to transfer light between two waveguides, the phase matching condition must be perfectly satisfied, or else the maximum transferable power will decrease significantly, and the coupling length will also vary. Therefore, it is not suited for hybrid structures, which are based multiple materials.

Adiabatic coupling rely on the supermode transformation described **Figure I-16(a)** and **(c)**. It can be seen that the even (odd) supermode will be mostly confined in the waveguide with the highest (lowest) propagation constant. Thus, by controlling the phase mismatch between the two waveguides, the power in either supermode can be transferred from one waveguide to the other. The most straight-forward way to control the propagation constants is by varying the width of the waveguides along the propagation direction, thus to form tapers. Compared to the directional coupling scheme, a perfect 0-100% power transfer is not achievable, and the length

necessary will also be longer [80]. Nevertheless, adiabatic coupling is also way more robust to achieve coupling, and thus more suited for hybrid structures.

In the case of the hybrid structure, only the supermode with the best trade-off between confinement in the active region and feed-back will lase. Generally, the even supermode is promoted in the designs, as shown on **Figure I-17**. Thus, the coupling length must be long enough to transform the supermode, without scattering power in the other modes (hence the name “adiabatic”). The taper (or mode transformer) design is treated later in **section III-1.2**.

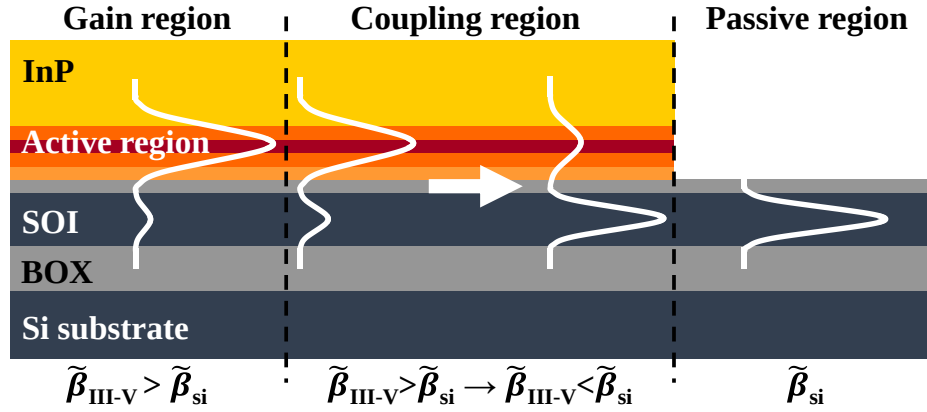


Figure I-17: Schematic view of the power transfer between the active and passive waveguides using adiabatic coupling.

Alternative structures

While most hybrid laser structures are similar to the one depicted on **Figure I-15**, alternative designs can also be found in the literature. Creazzo et al. proposed to etch a pit receptor in the SOI and BOX layers to metal-bond the III-V active region on the silicon wafer [66]. In this configuration, shown on **Figure I-18(a)**, the light is directly edge coupled into a thick ($1.5\mu\text{m}$) SOI waveguide. Since the metal layers used for the bonding are far from the active gain region, optical absorption is not an issue, and the direct connection with the substrate provides an excellent heat sink to evacuate the heat. Unlike the edge-coupled lasers presented in **section I-2.3**, the III-V gain region is not processed beforehand, and the waveguides are formed after bonding on the SOI. Thus, there are no misalignments between the III-V and silicon waveguide. However, the main issue of these lasers is the bonding of the III-V chip in an etched cavity, which alignment tolerances depend on the pit area. Thus, a pit much larger than the gain section is generally needed, and the gap between III-V and silicon waveguides is filled with SiO_2 and amorphous silicon [67].

Another design proposed by Matsuo et al. is to bond III-V materials as usual, to etch the gain region in a ridge waveguide, and then to use epitaxial growth on the sides of the gain region to form contacts regions [81], [82]. This way, the current is injected laterally, the electrodes can be formed far from the active gain region and thin lasers (below 300nm) can be formed. Up to now, this structure has been demonstrated has a stand-alone device, but has not been integrated with a SOI waveguide.

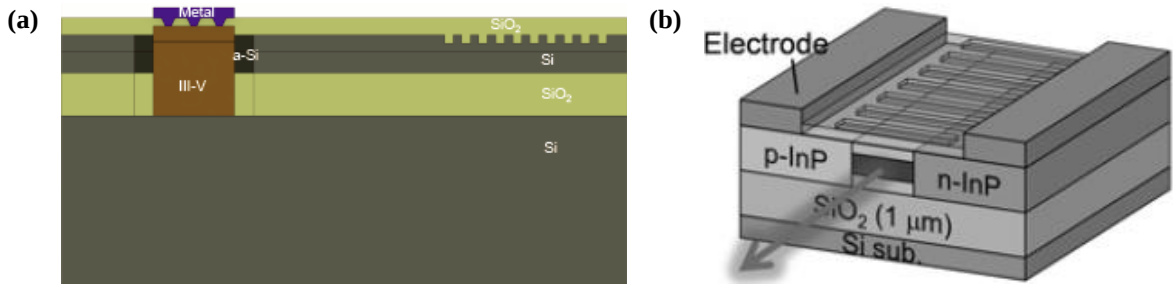


Figure I-18: (a) III-V waveguide metal bonded to the silicon substrate, edge coupled to a SOI waveguide [68]. (b) Active waveguide bonded on a silicon substrate, with re-grown contact regions after bonding for lateral injection [82].

I-3.3. Existing hybrid III-V on silicon laser cavity designs and state-of-the-art

In the last decade, several types of hybrid III-V on silicon lasers, with several feedback and mode-selection schemes have been demonstrated. The aim of this section is to list the different designs for laser cavities, and to quickly present their advantages and inconveniences. At the end of this section, the performances of the state-of-the-art of each cavity design are listed **Table I-1**.

Fabry-Pérot

The Fabry-Pérot (FP) cavity is the most straight-forward design. Feedback is provided by two mirrors reflecting over a large wavelength span, without any kind of mode-selection filter. Given their length (several hundreds of microns), a lot of wavelengths can fulfil the phase condition for lasing. Thus, FP lasers do not operate in single-wavelength regime, as illustrated by their low *SMSRs* in **Table I-1**. The main interest of F-P cavities is their simplified design, and they are often used to validate the hybrid structure. Thus, the first electrically-pumped demonstration of a hybrid III-V on silicon laser in continuous regime used a FP cavity [42]. The mirrors can be defined either by cleaved facets in the silicon waveguide [44], [59] (as shown on **Figure I-19**), a silicon waveguide fully-etched at its ends [45], or fully-etched silicon Bragg gratings [46]. The operating principles of Bragg gratings will be described later in **section III-1.3**.



Figure I-19: Schematic view of a Fabry-Pérot cavity with cleaved facets in the silicon waveguide. Image from [62].

Racetrack

Compared to FP cavities, racetracks are even easier to fabricate. Indeed, the feedback is simply provided by a silicon waveguide, used to link the two extremities of the hybrid waveguide, as displayed on **Figure I-20**. Thus, there is no need to fabricate mirrors through cleaved facets or full silicon etching. As in the case of FP lasers, the demonstrated hybrid III-V on silicon racetrack lasers generally do not have mode-selection filters. Since the cavity formed by a racetrack is long (1-2mm), a lot of wavelengths can satisfy the phase condition for lasing, thus their *SMSRs* are quite small. Due to their relatively low interest, there have been few demonstrations of racetrack cavities [46], [47].

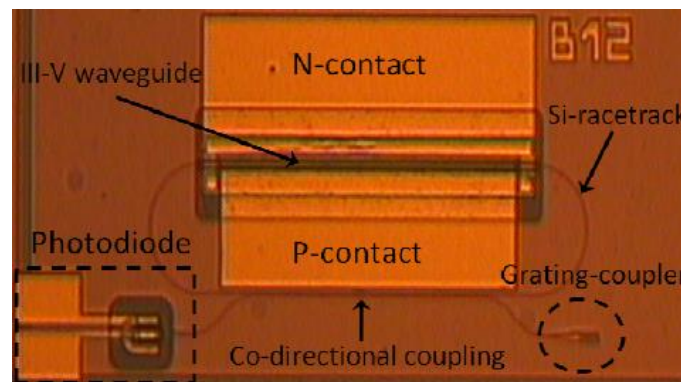


Figure I-20: Top microscopic view of a racetrack cavity. Image from [46].

Distributed Bragg reflector

In the case of a hybrid III-V on silicon distributed Bragg reflector (DBR) cavity, feedback is provided by two Bragg gratings located in the silicon waveguide, at each side of the gain region. Unlike FP cavities, the DBR lasers operate in the single wavelength regime. Indeed, shallowly-etched gratings also offer wavelength-selection

properties, as it will be discussed in **section III-1.3**. This type of laser has been demonstrated several times [48], [51], with better performances than FP lasers – in terms of current threshold, slope efficiency, and *SMSR* –, as shown in **Table I-1**. The main issue with these lasers is their total length (including the gratings), which is generally on the order of 1-2mm, as shown on **Figure I-21(a)**. To solve this issue, a solution is to use rings as wavelengths filters with the gratings, as depicted on **Figure I-21(b)**. This way, the gratings are only used as reflectors, and can be more deeply etched to reduce their lengths [52].

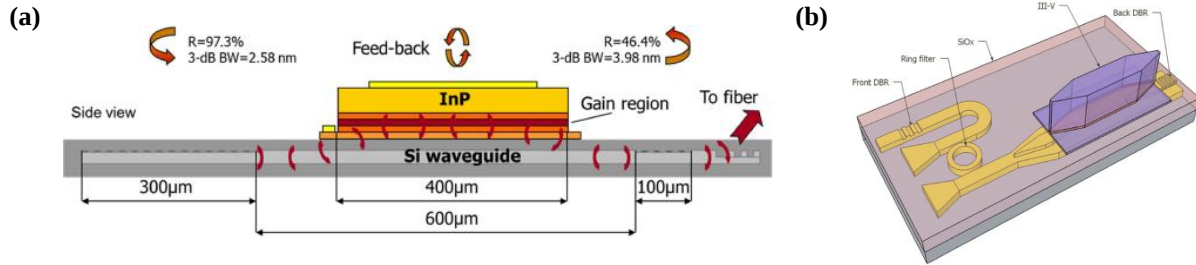


Figure I-21: (a) DBR laser with shallowly-etched gratings [57]. (b) Ring-assisted DBR laser [52].

DBR lasers also have an interesting feature, which is the possibility to tune the emitted wavelength. By changing the effective index of refraction of the wavelength filter, the phase condition of the cavity can be controlled. As it will be explained in **section II-1.3**, a local change of refractive index can be easily realised through the thermo-optical effect, by using a resistive layer as a heater, deposited above the wavelength filter. Wavelength tuning up to 20nm was demonstrated this way [51]. The tuning range can be even increased further, using specific designs such as sampled-grating DBR (SGDBR) lasers, based on the Vernier effect [53], [67].

Distributed feedback

Hybrid III-V on silicon distributed feedback (DFB) lasers also use a Bragg grating, but instead of being outside of the gain region, it is located below. Thus, they have a smaller footprint than DBR cavities. While it is possible to form a laser with a single grating, a short section between two gratings is generally located near the center of the gain region, in order to ensure single-wavelength operation. Even if the length of this section can vary according to the designs [58], single-wavelength DFBs generally use a section as long as a grating period, also referred as a “quarter-wavelength phase-shifted” section [49]–[51], [60]. An example of DFB laser with a quarter-wavelength shifted section is displayed on **Figure I-22**. If the DFB is symmetric, it will emit the same power in both output directions. The operating principle of the DFB laser based on a quarter wavelength phase-shifted section will be described in **section IV-2.1**.

As it can be seen on **Table I-1**, state-of-the-art hybrid III-V on silicon DFB lasers have shown excellent performances in terms of current threshold, maximum power outside the cavity, slope efficiency and *SMSR*. The main issue with DFB lasers is that since the silicon waveguide is right below the gain region, its refractive index of refraction cannot be easily controlled through heating devices as for DBR lasers. Thus, DFB lasers with wavelength tuning capabilities have not been demonstrated up to now.

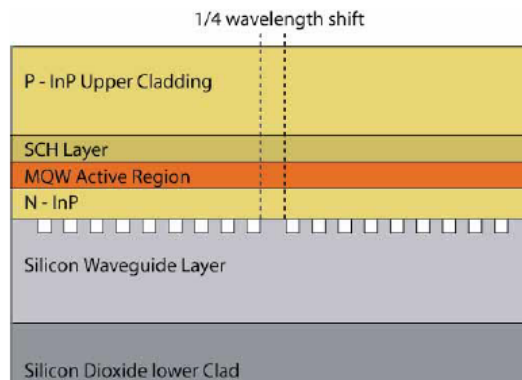


Figure I-22: Schematic view of a DFB laser with a quarter-wavelength shift. Image from [43].

Microdisks and microrings

Microdisk cavities were amongst the first demonstrated with electrical pumping in continuous regime for hybrid III-V on silicon lasers, shortly after the F-P cavities [54]. As for the racetrack cavities, the feedback is not provided by a mirror. Instead, a whispering gallery mode circulates along the edge of the III-V microdisk, and a fraction of the generated optical power is evanescently coupled in a silicon waveguide. An example of hybrid microdisk structure is shown on **Figure I-23(a)**. Given small the diameter of the microdisks ($< 10\mu m$), the length of the cavity is quite small. Thus, a reduced number of wavelengths will fulfil the phase condition for lasing, and a wavelength-selection filter is not mandatory to reach single-wavelength operation. As it can be seen on **Table I-1**, while their output power is quite low compared to the other types of cavities, their reduced current threshold and footprints make them ideal candidates for short communications distances, such as intra-chip connections.

Another type of circular cavity has also recently gained a renewed interest for these applications: the hybrid III-V on silicon microring laser [55]. In this case, a silicon microdisk is present below the III-V microring, as displayed on **Figure I-23(b)**. The optical power is confined in the silicon waveguide, and evanescently coupled in the III-V gain region. While being quite small (with a diameter $< 50\mu m$), they generally remain larger than microdisks cavities. Thus, their current thresholds are larger than for the microdisk, but their output powers are also higher. The main issues with these structures are their difficulty of fabrication due to their small diameters, as well as their high series resistance, which limits their input electrical power and their performances at higher temperatures.

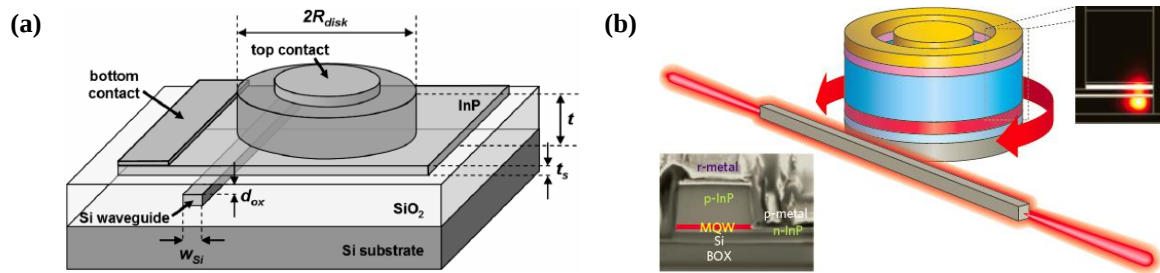


Figure I-23: Schematic views of hybrid III-V on silicon (a) microdisk [75], and (b) microring [15] lasers.

Table I-1: State-of-the art of the existing heterogeneously integrated III-V on silicon laser structures.

All presented results are in continuous regime, at ambient temperatures as close as possible to 25°C.

Structure characteristics				Figures of merit							Source		
Cavity structure	Mode-selection filter	Bonding scheme	Coupling type	J_{th} (kA/cm ²)	I_{th} (mA)	P_{max} (mW)	Slope efficiency (mW/mA)	λ_p (nm)	$SMSR$ (dB)	Wavelength tuning (nm)	Laboratory or Company	Year	Ref.
Fabry-Pérot	none	Direct	Evanescent (confined in Si)	Unknown	35 ^a	24.2 ^{a,b}	Unknown	1326	Unknown	none	UCSB/Intel	2007	[44]
		Direct	Evanescent (confined in Si)	1.25	60	12.5	0.07	1490	< 10 ^a		Caltech	2009	[45]
		Direct	Adiabatic	0.9	50	13 ^a	Unknown	1588	12 ^a		CEA-LETI	2012	[46]
		Adhesive	Adiabatic	Unknown	30	4.5	0.043	1590	< 5 ^a		IEF/IMEC/III-V Lab	2012	[59]
Racetrack	none	Direct	Adiabatic	Unknown	30	2.8	Unknown	1550	Unknown	none	CEA-LETI	2012	[46]
		Direct	Evanescent (confined in Si)	Unknown	200 ^a	23 ^{a,b}	Unknown	1592	7 ^a		UCSB/Intel	2007	[47]
Distributed Bragg reflector	Grating	Direct	Evanescent (confined in Si)	Unknown	65	8	Unknown	1598	50	none	UCSB/Intel	2008	[48]
	Grating	Direct	Adiabatic	Unknown	17	15	0.107	1547	52	20	CEA-LETI	2013	[51]
	Ring	Direct	Adiabatic	Unknown	38	4	Unknown	1553	45	8	IMEC/ III-V Lab	2013	[52]
	SGDBR	Direct	Adiabatic	Unknown	83	7.5	Unknown	1300	35	35	CEA-LETI	2016	[53]
	SGDBR	Metallic	Edge-coupling	Unknown	75	30	0.125	1560	45	Unknown	Skorpios	2014	[67]
Distributed Feedback	Quarter-wavelength phase-shifted section	Direct	Evanescent (confined in Si)	1.1	8.8	3.7 ^b	Unknown	1575	55	none	UCSB	2014	[49]
		Direct	Adiabatic	1.03	36	22 ^b	0.24	1305	55		CEA-LETI	2015	[50]
		Direct	Adiabatic	Unknown	65	40 ^b	Unknown	1560	> 30		CEA-LETI	2014	[51]
		Adhesive	Adiabatic	1.72	35	14	0.135	1566	50		IMEC/ III-V Lab	2013	[60]
Microdisk	none	Direct	Evanescent (confined in III-V)	1.13	0.5	0.01	0.03	1600	< 30 ^a	none	IMEC/INL/CEA-LETI	2007	[54]
Microring	none	Direct	Evanescent (confined in Si)	2.9	12	0.2 ^b	Unknown	1529	< 15 ^a	none	UCSB/HP labs	2009	[55]

^a Estimated from graphical data.^b Total power emitted at both outputs of the cavity.

I-3.4. Advantages of hybrid III-V on silicon lasers and their integration issues

Although they are not monolithically integrated, hybrid III-V on silicon lasers present several advantages compared to externally integrated laser sources. First of all, even if there are some optical losses during the coupling between III-V and silicon waveguides, there are no additional optical losses between the laser cavity output and the input of the silicon photonic circuit. Besides, tapers used for adiabatic coupling have already demonstrated losses below 0.6dB [83]. Secondly, since the III-V waveguides are formed by photolithography after being bonded on the SOI wafer, there will be a reduced misalignment between the III-V and silicon waveguides, and an accurate alignment is not required during the bonding of III-V materials, even in the case of die bonding. Thirdly, since the laser sources are integrated at the wafer-level, the packaging of the silicon photonic circuits will be greatly simplified, and its overall cost reduced. Finally, the heterogeneous integration of III-V materials on silicon offers new possibilities in the design of the lasers. Even though silicon is not a good light emitter, all the other functions necessary for the laser (feedback, wavelength selection, and output coupling) can be made in silicon. While those functions can also be realized in the III-V waveguide, it is way more relevant to use silicon to realize the maximum of these functions in order to make use of the mature and well-controlled fabrication processes of the CMOS fabrication lines. This way, reliable and highly-performant silicon components can be fabricated with high yields, while leaving the optical gain function to the III-V materials.

Nevertheless, even with all these advantaging features, the hybrid III-V on silicon lasers demonstrations found in the literature are generally for stand-alone components or at best associated with passive silicon components such as waveguide-to-fiber grating couplers. The reasons for this lack of more complex silicon photonic circuits with integrated laser sources are mainly due to the integration of III-V materials on a silicon photonic platform. While it was explained in **section I-2.3** that III-V materials have received a renewed interest from the microelectronics industry, they are still not considered as standard materials in the CMOS fabrication platforms, and thus not fully used for silicon photonics. Another issue comes from the metals used to form the electrical contacts on the III-V waveguide. Those contacts are generally gold-based, since it offers the best contact resistivity with the III-V materials. However, gold is prohibited in CMOS fabrication platforms since it is a dopant for silicon. Additionally, as it will be detailed in **section III-1.1**, the silicon waveguide must be thick enough in order to efficiently transfer the light from the III-V active waveguide. Therefore, most demonstrated hybrid III-V on silicon lasers rely on silicon waveguide with a thickness larger than 400nm . This is an issue knowing that most of the previously cited silicon photonics fabrication platforms use starting SOI layers with thicknesses below 300nm . Starting from a thicker SOI layer would be equivalent to re-define the other silicon components design and fabrication steps, which is an issue for mature fabrication platforms. Finally, as it will be shown in **section III-1.1**, a typical III-V stack used for hybrid III-V on silicon lasers is 2 to $3\text{-}\mu\text{m}$ -thick. After bonding, the large topography induced by the III-V stack will prevent the use of conventional multi-level metallization, which have approximately the same height. If the III-V lasers are bonded first, the standard multi-level metallization which needs a planar surface between each level will be impossible to process. And if the III-V laser is bonded after the metallization, the optical power transfer via adiabatic coupling will become impossible due to the large gap between III-V and silicon waveguides.

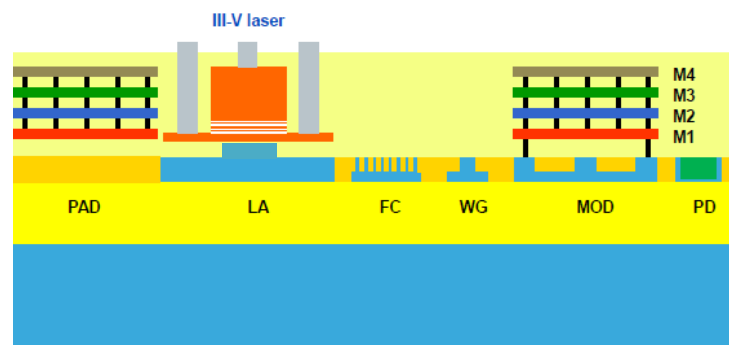


Figure I-24: Schematic view of a silicon photonic circuit comprising different kinds of passive and active devices, and illustrating the issues of laser integration in terms of SOI layer thickness and standard multi-level metallizations. This co-integration is not possible in these conditions.

To summarize, in its current state, the laser integration will disturb the fabrication of the other components which have been developed beforehand. This is why it is not widespread in complex silicon photonic circuits. That being said, demonstrations of high-speed silicon photonics transmitters using a hybrid III-V on silicon laser have been realized and are discussed in the next section.

I-4. High-speed integrated transmitters

While the integration of hybrid III-V on silicon laser to form a high-speed integrated transmitter is not a trivial task, a few teams have taken the challenge and proposed some prototypes. However those rare demonstrations also had their limitations in terms of performances. Therefore, during this PhD thesis, our aim was to fill this void, and to study how to provide a high-speed transmitter with an integrated laser for silicon photonics. In this section, the current state-of-the-art of integrated transmitters is first discussed, with their performances and limitations, followed by the objectives of this work.

I-4.1. State-of-the-art for integrated transmitters in silicon photonics

As explained in **section I-2.2**, one of the main objectives of silicon photonics is to provide high-speed transmission for the next generation of data-centers. To the best of our knowledge, there are only two groups who demonstrated the co-integration of hybrid III-V on silicon lasers and silicon modulators to form high-speed transmitters. Alduino et al. have shown four hybrid III-V on silicon DBR lasers, each emitting a different wavelength fixed in the $1.3\mu\text{m}$ region, co-integrated with four silicon modulators, each operating at a 12.5Gb/s data rate, thus leading to a total modulation rate of 50Gb/s [84], [85]. A schematic view of their transmitter is shown on **Figure I-25(a)**. Very recently (August 2016), Intel announced starting the shipments of $4 \times 25\text{Gb/s}$ modules supposedly based on the same technology, but without providing specific information on their characteristics. Duan et al. have also demonstrated a similar transmitter, composed of a single but tunable hybrid III-V on silicon DBR laser, emitting in the $1.55\mu\text{m}$ region, and a silicon modulator limited at a 10Gb/s modulation rate [86]. The maximum wavelength tuning was 9nm . A schematic view of their transmitter is displayed on **Figure I-25(b)**. While providing less modulation rate than the other demonstration, the wavelength tuning is an essential feature to have fine wavelength spacing between the different channels of the transmitter, as it will be explained in **section III-1.3**.

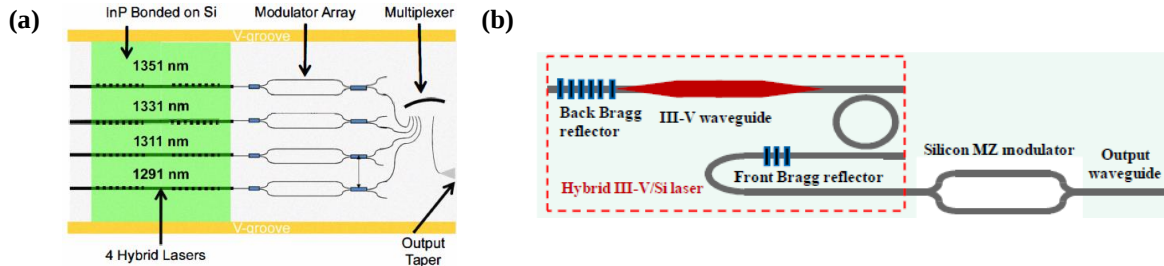


Figure I-25: Schematic views of high-speed transmitters co-integrating hybrid III-V on silicon lasers and silicon modulators, with : (a) four lasers with a fixed wavelength [85], and (b) a single tunable laser [86].

Another possible approach for a transmitter based on the hybrid lasers would be to not only use the III-V heterostructures brought by bonding to form the lasers, but also the other active components, such as modulators and photodetectors. In this case, the SOI layer is only used to form the passive components, and light is transferred several times between III-V and silicon waveguides. With this solution, some of the integration issues listed in **section I-3.4** could also be avoided, since the III-V regions would be the only ones with metallic contacts. An example of transmitter based on this approach is shown on **Figure I-26**. By using this approach, Ramaswamy et al. demonstrated four hybrid III-V on silicon tunable lasers, each emitting a different wavelength fixed in the $1.3\mu\text{m}$ region, co-integrated with four III-V on silicon modulators, each operating at a 28Gb/s data rate [87]. The advantages and limitations of this type of transmitter will be further discussed in **section II-1.3**.

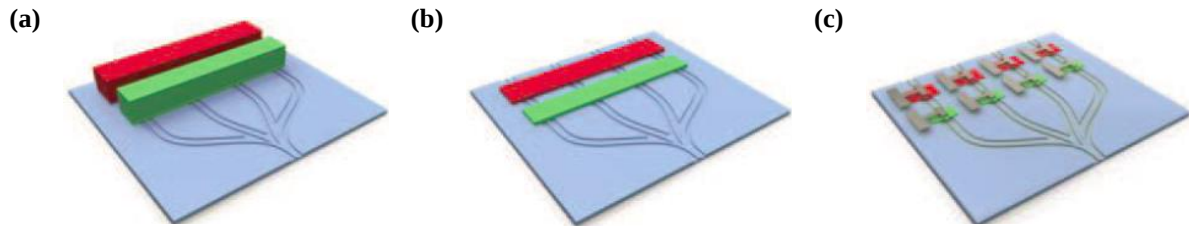


Figure I-26: Hybrid III-V on silicon integrated transmitter, where the III-V is not only used to form the laser, but all active components (modulators and photodiodes). Schematic views of the transmitter (a) after bonding, (b) after substrate removal, and (c) after etching and metallization processes. Images from [88].

Finally, even if in **section I-2.2** the transmitter presented as an example was based on an indirect modulation scheme – with a laser and a modulator –, the direct modulation of the laser source is another possibility which can also be considered. Recently, a directly-modulated hybrid III-V on silicon DFB laser has been shown with a modulation rate of 28Gb/s , which is the highest demonstrated up to now [89]. This solution is more compact, but a directly modulated laser will also offer less extinction ratio and electro-optical bandwidth (both defined in **section II-1.2**) than in an indirect modulation scheme [90]. Additionally, when the input current of the laser is modulated, its gain and its refractive index will both be modulated, resulting in a modulation of both output power and resonating wavelength (or frequency). If a modulation of intensity is desired, the frequency modulation (or chirp) will broaden the modulated spectrum of the laser and add dispersion penalties, which increases with the distance and the modulation rate [3]. On the other hand, chirp is smaller (and can even be close to 0) in indirect modulation schemes. Finally, if the laser is not directly modulated, a trade-off will not be necessary between its static and dynamic performances. Thus, in the rest of the document, we will focus our attention on indirect modulation schemes.

I-4.2. Work objectives and targeted specifications

In this chapter, we have seen that, even with its numerous qualities, silicon photonics must still face a major hurdle, which is the lack of monolithically integrated laser sources. Even if the optical source can be externally integrated, a wafer-scale integrated laser source would bring several advantages in terms of coupling losses, tolerance towards spatial misalignments and packaging simplicity. While the heterogeneous integration of III-V material remains a promising solution in order to obtain optical gain on silicon, demonstrations of fully integrated transmitters are still not widespread. This is mainly due to the complex co-integration of III-V lasers with other active components, which is not a trivial task. Therefore, the main goal of this PhD thesis is study the conception of a hybrid III-V on silicon integrated transmitter for silicon photonics. This transmitter should be as compatible as possible with current silicon fabrication platforms, which will lead to several trade-offs during the design of the components.

As it was explained in **section I-2.1**, one of the current aims of silicon photonics is to propose data transmission solutions for the next generation of data-centers based on SMF transmission, for distances up to 10km . Therefore, we decided to target the same application for the transmitter. In order to choose the targeted performances of the transmitter, the ideal way would be to associate it to a defined receiver. Thus, the sensitivity of the receiver would have fixed the performances needed for the transmitter, in order to realize a high-speed transmission link. Since we did not have a specific receiver in mind, we decided to use some of the specifications of the 100GBASE-LR4 norm (from the IEEE-802.3ba standard) as a reference for the targeted performances of the transmitter. Therefore the conceived transmitters should be able to transmit information up to 10km , at four different wavelength in the $1.3\mu\text{m}$ wavelength region. The wavelength spacing is also fixed at 4.5nm . The lasers used in the transmitters must operate in the single-wavelength regime, and each wavelength or channel must be modulated at a 25Gb/s bit rate. Finally, the optical modulation amplitude (which is a specific FOM used to describe the overall modulation performances of a transmitter, and will be defined in **section II-1.2**) of each channel should be above -1.3dBm . This targeted limit will be used as a basis for discussion, since a receiver with an excellent sensitivity might detect without any errors the transmitted information for lower optical modulation amplitude values. The targeted performances for the transmitters are all summarized in **Table I-2**.

As it can be seen on **Table I-2**, some of these requirements depends solely the laser (wavelength, *SMSR*), or on the modulation system (modulation rate). However, the transmission distance and the optical modulation amplitude depend on both the laser output power and on the modulator optical losses and modulation efficiency, and are tougher to predict if the performances of one of the two components is not known beforehand. Given the wavelength spacing between the channels, a tunable laser will also be necessary for the transmitter.

Table I-2: Targeted performances for the transmitter.

Description	Target	Unit
Distance (max)	10	km
Wavelengths (or channels)	4 in the O-band	
Wavelength spacing	4.5	nm
<i>SMSR</i> (min)	30	dB
Modulation rate per channel	25	Gb/s
Optical modulation amplitude per channel (min)	-1.3	dBm

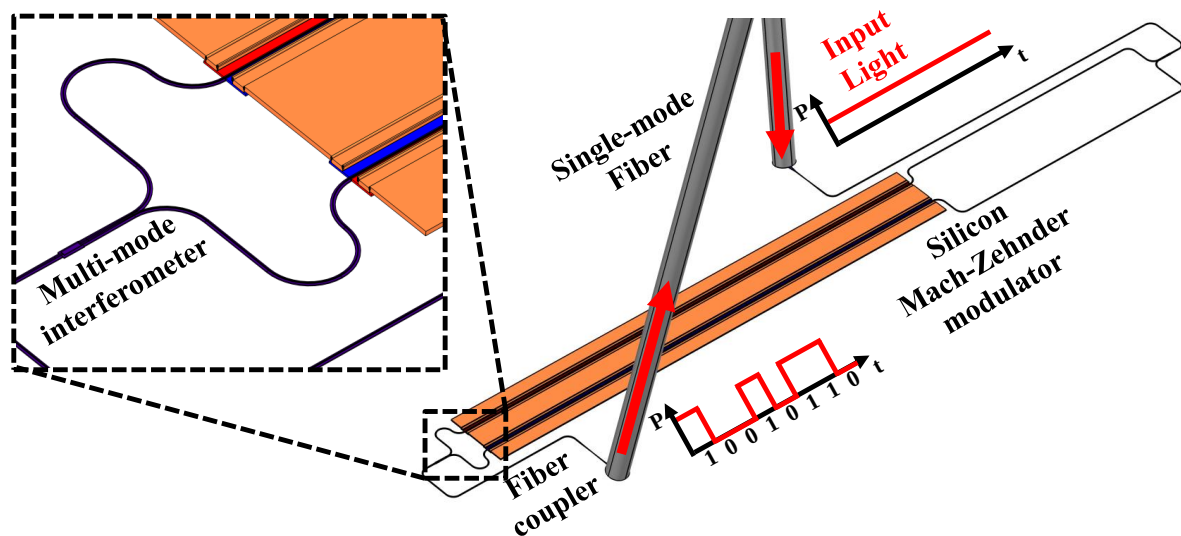
Another objective of this work is to improve the integration of the hybrid III-V on silicon lasers with the rest of the silicon photonic circuit. The search of CMOS-friendly metallic contact for the laser, as well as the utilization of standard multi-level metallization even with the large on-wafer topography induced by the thick III-V waveguide are not tackled in this work, but are the subject of others PhD theses (by Elodie Ghégin and Jocelyn Durel). Nevertheless, the silicon thickness difference issue between what is needed for the laser, and what is currently proposed in the standard silicon photonic platforms is effectively treated in this work.

The document is organized in the following order. In **chapter II**, a stand-alone modulator, designed to be co-integrated with a hybrid III-V on silicon laser source is presented. The reasons which lead to the choice of the modulator structure, as well as its conception, fabrication and characterization are also detailed. In **chapter III**, the design of the laser source used for the integrated transmitter is first presented, followed by the fabrication and characterization of the complete transmitter. The solution proposed to solve the SOI thickness difference between the laser and the other components is presented in **chapter IV**. This solution was not directly implemented on the complete transmitter, but was used for the demonstration of a laser associated with a surface-to-waveguide grating coupler. Finally, conclusions on this work, as well as future perspectives are presented in **chapter V**.

It can be noted that this PhD thesis was realized in parallel of the one of H     Duprez (who successfully obtained her PhD earlier this year), whose work was focused on the design of hybrid III-V on silicon lasers. Even if we were both trained in the design of these optical sources, she pursued the optimization of the lasers used in this work, while I was working on the modulators, so that both components could be ready at the same time for their co-integration. Given the large amount of time necessary to fabricate the transmitters, it would have not been possible to obtain a complete device otherwise. As for the design of the modulators, I received a large amount of help from Dr. Benjamin Blam     (from CEA-LETI), who initiated me to the world of RF propagation and shared his knowledge on the subject during this thesis. Nevertheless, I realized all the analyses and simulations presented in this work, by adapting to my needs pre-existing models used for the lasers and travelling-wave electrodes, and also by developing models myself for the modulators active region and eye diagrams simulation. While I did not operate directly the various equipment used during the fabrication, I spent a lot of time in the 8'' cleanroom of CEA-LETI in order to follow the fabrication processes of the various components. I interacted as much as possible with the technicians and process engineers, in order to make the best possible decisions with the information available at a given time. I also realized or assisted to almost all the various characterizations performed during the fabrication (such as profilometry, ellipsometry, scanning electron microscopy), those information being essential to comprehend the results of the terminated devices characterization. Even if I could not follow as precisely the III-V processing steps which were realized in different cleanroom, I shared as much information as possible with Dr. Christophe Jany and his team, in order to fully understand the devices I would be testing. I also realized the various electro-optical characterizations of the terminated devices, and interacted with the characterization team from CEA-LETI, who taught me all I needed to know to do my measurements, and assisted me in time of needs. Finally, I performed the complete analysis of the characterizations and tried as much as possible to link them to the simulations and the processes used during the fabrication.

Chapter II: Silicon Mach-Zehnder modulator

This second chapter is focused on the developments of silicon Mach-Zehnder modulators (MZMs) designed to be co-integrated with hybrid III-V on silicon lasers, thus forming the hybrid III-V on silicon high-speed transmitters. First, a general study of the optical modulation on silicon is presented, as well as the existing structures for modulation, in order to justify the choice of using MZMs. Then, the design of the MZMs, which include p - n junctions and travelling-wave electrodes (TWE), is thoroughly detailed. In order to validate the design of the modulators, standalone devices have been fabricated and tested before the complete transmitters. The fabrications steps – later used for the transmitter – as well as the results of the measurements are both described in this chapter. Finally, solutions are proposed to improve further the performances of the modulators.



II-1. Design principles and state-of-the-art.....	26
II-1.1. General principles of optical modulation	26
II-1.2. Figures of merit.....	27
II-1.3. Physical phenomena for optical modulation in silicon	29
II-1.4. Device structures used for carrier depletion devices	33
II-1.5. Target performance of the modulator	35
II-2. Silicon Mach-Zehnder modulator design	35
II-2.1. Junction design	35
II-2.2. Speed limitations and travelling-wave electrodes design	39
II-2.3. Device architecture.....	46
II-3. Silicon Mach-Zehnder modulator fabrication	47
II-4. Mach-Zehnder modulator characterization.....	51
II-4.1. Passive components optical characterization	51
II-4.2. Active region static electro-optical characterization	52
II-4.3. Small-signal electrical and electro-optical characterization	53
II-4.4. Large-signal electro-optical characterization	56
II-5. Conclusions and possible improvements	62

II-1. Design principles and state-of-the-art

Before considering the design of a silicon modulator, it is essential to know the existing solutions for optical modulation. Indeed, there are several devices, which are more or less suited depending on the targeted system characteristics. The aim of this section is to briefly present the possible choices for the modulator, and explain why the silicon Mach-Zehnder based on carrier depletion was chosen for our application. To start the discussion, it is important to define the basis of optical modulation and where to focus our attention.

II-1.1. General principles of optical modulation

In order to transmit information using photonic systems, a component is required to encode information from the electrical realm into the optical one. Unless the laser is directly modulated, this is the role of the optical modulator, which will convert an input optical signal with constant power, into a signal with variations of its amplitude (intensity modulation) or its phase (phase modulation). In this document, only intensity modulation is discussed.

In order to encode information on a single channel (or a single wavelength), several kinds of intensity modulation exist. Amongst them, the On-Off Keying (OOK) is the simplest form. In OOK, the most generally used modulation format is non-return-to-zero (NRZ), where two levels of power are used to differentiate two states of information: an “ON” power level (P_{ON}), equivalent to a bit value of 1, and an “OFF” power level (P_{OFF}), equivalent to a bit value of 0. By alternating those power levels for a given bit rate (or bit duration T_{bit}), it is possible to transmit information. An example of optical modulation based on NRZ format is presented on **Figure II-1**, where an electrical signal is used to control the modulator, and transfer the information in the optical domain.

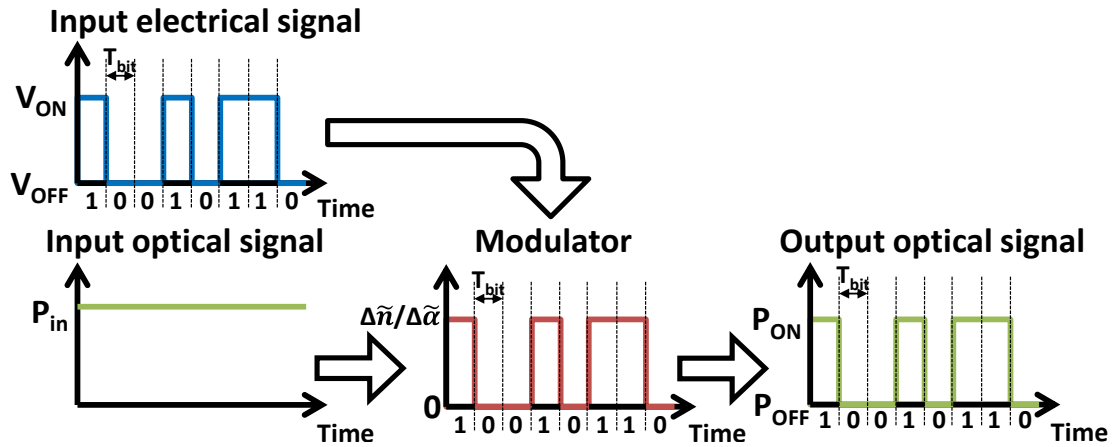


Figure II-1: General principle of optical intensity modulation.

The general principle of optical modulation can be explained as the following. First, we will consider the case of light propagating in a waveguide fundamental TE mode, along the z -axis, which electric field ($\vec{E}$) is expressed by **Eq. (I.3)**. Since there is no optical gain expected in this waveguide, the complex propagation constant ($\tilde{\beta}$) is defined by:

$$\tilde{\beta} = \frac{2\pi\tilde{n}}{\lambda} - j\frac{\tilde{\alpha}}{2} \quad (\text{II.1})$$

Where $\tilde{n}$ is the mode effective index of refraction and $\tilde{\alpha}$ is the transverse modal loss. Both of these coefficients depend on the spatial distribution of the refractive index $n(x, y, z)$ and absorption coefficient $\alpha(x, y, z)$. By combining both **Eqs. (I.3)** and **(II.1)**, we get:

$$\vec{E}(x, y, z, t) = \vec{e}E_0U(x, y)e^{j\omega_0t}e^{-j\left(\frac{2\pi\tilde{n}}{\lambda} - j\frac{\tilde{\alpha}}{2}\right)z} = \vec{E}(x, y, t)e^{-\frac{\tilde{\alpha}}{2}z}e^{-j\frac{2\pi\tilde{n}}{\lambda}z} \quad (\text{II.2})$$

Since the optical intensity is proportional to the magnitude squared of the electric field, it can be seen using equation **Eq. (II.2)** that intensity modulation is directly obtained by creating a variation on the transverse modal

loss ($\Delta\tilde{\alpha}$). Devices based on this principle are known as electroabsorption modulators (EAMs). However, intensity modulation can also be reached with a variation on the effective index of refraction ($\Delta\tilde{n}$) – thus the phase of the electrical field – by using a resonant, or an interferometric structure. Devices based on this principle are referred as electrooptic modulators (EOMs).

II-1.2. Figures of merit

In order to evaluate the performances of the modulators, several FOM are used. These FOM describe the efficiency, losses and speed of the modulators. Generally, they cannot be optimized individually, and a trade-off must be found between them.

Efficiency

The first FOM is the extinction ratio (ER), which express the ratio between the “ON” and “OFF” states output powers of the modulator, and which must be maximized to avoid errors during the transmission of data:

$$ER = \frac{P_{ON} [mW]}{P_{OFF} [mW]} \geq 1 \quad \text{or} \quad ER_{dB} = 10 \log_{10} \left(\frac{P_{ON} [mW]}{P_{OFF} [mW]} \right) \geq 0 \quad (\text{II.3})$$

In the case of EOMs, which are based on phase variation, another FOM is also considered for the efficiency, which is the linear phase shift ($\Delta\varphi$) for an applied voltage (V_{bias}):

$$\begin{aligned} \Delta\varphi (V_{bias}) [rad/mm] &= \frac{2\pi\tilde{n}(V_{bias})}{\lambda} - \frac{2\pi\tilde{n}(0)}{\lambda} = \frac{2\pi\Delta\tilde{n}(V_{bias})}{\lambda} \\ \text{or} \quad \Delta\varphi (V_{bias}) [^\circ/mm] &= \frac{360\Delta\tilde{n}(V_{bias})}{\lambda} \end{aligned} \quad (\text{II.4})$$

Another FOM often used to describe the efficiency of EOMs is the $V_\pi L_\pi$ product. For a given voltage (or linear phase shift), L_π is defined as the modulator length needed to reach a total phase shift (defined by $\Delta\varphi \times L_m$, with L_m the modulation length) equal to π . Therefore:

$$V_\pi L_\pi (V_{bias}) = V_{bias} \frac{\lambda}{2\Delta\tilde{n}(V_{bias})} \quad (\text{II.5})$$

This FOM is quite useful, especially if $\tilde{n}$ is directly proportional to V_{bias} . In this case, the $V_\pi L_\pi$ product becomes constant for any voltage, and gives direct information on the modulator length needed for modulation. Even if it is not the case, this FOM can still be used, and easily converted in linear phase shift using Eq. (II.6):

$$\begin{aligned} \Delta\varphi (V_{bias}) [rad/mm] &= \frac{\pi * V_{bias}}{10 * V_\pi L_\pi (V_{bias}) [V.cm]} \\ \text{or} \quad \Delta\varphi (V_{bias}) [^\circ/mm] &= \frac{18 * V_{bias}}{V_\pi L_\pi (V_{bias}) [V.cm]} \end{aligned} \quad (\text{II.6})$$

Both the linear phase shift and $V_\pi L_\pi$ product are related to the extinction ratio, depending on the chosen EOM structure. Thus, the ER increases with an increasing $\Delta\varphi$ or a decreasing $V_\pi L_\pi$.

Losses

Even with a large extinction ratio, if the output optical power is too low compared to the receiver sensitivity, it will results in a noisy signal, and the ON and OFF states will not be distinguished, resulting in data transmission error. Therefore, the overall optical losses coming from the modulator must be reduced as much as possible in order to reduce the optical power needed to read the signals. These losses are generally defined by the insertion loss (IL) parameter, which must be as close as possible to 1 (or 0dB):

$$IL = \frac{P_{ON} [mW]}{P_{IN} [mW]} \leq 1 \quad \text{or} \quad IL_{dB} = \left| 10 \log_{10} \left(\frac{P_{ON} [mW]}{P_{IN} [mW]} \right) \right| \geq 0 \quad (\text{II.7})$$

Where P_{ON} is the output power at the ON state, and P_{IN} is the input optical power.

Speed

After the efficiency and losses, the last essential characteristic of a modulator is its speed. It can have intrinsic limitations due the physical phenomenon chosen for modulation, or due to a large capacitance for instance. It is estimated by the small-signal frequency response $M(\omega_e)$, which describes the behaviour of the component against the frequency of the input electrical signal (ω_e being the input electrical signal angular frequency). It is evaluated by sending a small-signal sinusoidal input voltage at angular frequency ω_e , and by observing the small-signal sinusoidal optical power at the modulator output (due to the linearity of the small-signal assumption). A static component can also be added to the input signal, and influence the modulator response. The frequency response is usually normalized with the low-frequency (or static) response. The normalized parameter $m(\omega_e)$ is called the modulation frequency response (or modulation depth), and is defined as [90]:

$$m(\omega_e) = \left| \frac{M(\omega_e)}{M(\omega_0)} \right| \quad \text{or} \quad m_{dB}(\omega_e) = 10 \log_{10} \left(\left| \frac{M(\omega_e)}{M(\omega_0)} \right| \right) \quad (\text{II.8})$$

with $\omega_e = 2\pi f_e$

Where f_e is the electrical signal frequency, and ω_0 is the lowest output angular frequency of the signal generator. From the modulation frequency response, an electro-optical (E/O) bandwidth ($f_{E/O}$) can also be defined. The whole frequency response is needed to understand the high-speed behaviour of the modulator, but this bandwidth is generally a good indication of the speed limitations of the modulator, and is defined as:

$$m(f_{E/O}) = \frac{1}{2} \quad \text{or} \quad m_{dB}(f_{E/O}) = -3 \text{ dB} \quad (\text{II.9})$$

Overall performances

Once those modulation parameters have been evaluated, the best way to know if the modulator can effectively transmit information is to study its dynamic large-signal response, and to control it with an actual NRZ electrical signal. To study all types of ON-OFF transitions, a pseudo-random binary sequence (PRBS) of bits is sent to the modulator. To be relevant, this sequence is typically between 2^7-1 and $2^{31}-1$ bit long, before repeating itself. In order to visualize those measurements efficiently, the usual representation is an eye diagram, where all the ON-OFF combinations of the PRBS sequence are synchronized and overlapped. Examples of eye diagrams are displayed on **Figure II-2**. These eye diagrams help to visualize the quality of the transmitted signal. A closed “eye” can be due to a low extinction ratio, a slow response of the modulator (leading to large ON-OFF transition times), or optical power levels below the sensitivity of the detector (due to the insertion losses). The last case is illustrated by **Figure II-2(a)**, where the optical power of the ON level is at the limit of the detector sensitivity, leading to a noisy measurement. In **Figure II-2(b)**, the modulator is used in the same driving conditions, but with an optical amplification system at its output to compensate for the optical losses, leading to a noise reduction, and a clearly opened eye. It also shows that the eye diagram largely depends on the set-up used during the measurement, and not only on the modulator. The eye diagram measurements permit to extract the extinction ratio and the insertion losses at a given data rate.

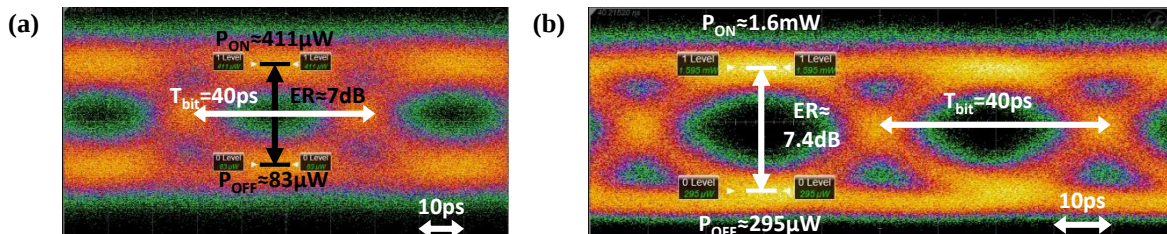


Figure II-2: Eye diagrams measurements at 25Gb/s of the same modulator, in the same driving conditions, but with different level of power measured by the detector: (a) low optical power (without output amplification), (b) high optical power (with output amplification).

Another parameter, the optical modulation amplitude (*OMA*) can also be used to describe the performances of the modulator. Defined in **Eq. (II.10)**, it can be seen that the *OMA* exhibit the advantage of including both

extinction ratio and insertion losses. High extinction ratio and low losses lead to a high *OMA*, and thus fewer errors in the signal transmission. It can also be expressed with the average power (P_{avg}) of the modulated signal. While it is convenient, it also depends on the modulator input optical power, which must be defined.

$$OMA [mW] = P_{ON} - P_{OFF} = P_{IN} * IL * \frac{ER - 1}{ER} = 2 * P_{avg} * \frac{ER - 1}{ER + 1} \quad (II.10)$$

$$P_{avg} = \frac{P_{ON} + P_{OFF}}{2} \quad (II.11)$$

Additional metrics

Another important metric is the optical bandwidth of the device – not to be confused with the electro-optical bandwidth $f_{E/O}$, which refers to useful operational wavelength range of a device, and can vary from below hundreds of picometers to over tens of nanometers. While the targeted optical bandwidth depends on the application, narrowband devices will be more sensitive to fabrication and temperature variations [91].

The footprint of the modulator can also be taken into account, since the length necessary for modulation can vary from tens of micrometers to several millimeters. While it might represent an issue in terms of component density, long devices will also suffer from large insertion losses and high-speed limitations (due to large capacitances) if not designed carefully [90].

II-1.3. Physical phenomena for optical modulation in silicon

Several effects can be used to induce directly (or indirectly) a variation on the effective index refractive index and/or the transverse modal loss in silicon-based structures using an electrical signal. Some are based intermediary physical effect (such as temperature) or make use of an additional material (such as Ge or InP). These effects are briefly described in this section. All of them affect the material refractive index and absorption spatial distributions, since they are related by the Kramers-Kronig relations [5]. Since each effect has different pros and cons, they must be studied and compared, in order to choose the most suited for our application. The study is focused on the $1.3\mu m$ wavelength region, which is the targeted wavelength of operation.

Pockels and Kerr effect

Optical modulation can be reached by directly applying an electric field on the material. Indeed, forces resulting from the electric field will distort the material crystalline structure, leading to a change of its refractive index [5]. If the refractive index changes are proportional to the electric field, the effect is referred as Pockels effect, and if the changes are proportional to the square of the electric field, the effect is known as Kerr effect. Pockels effect does not appear in centrosymmetric materials. Therefore it cannot be used directly for modulation in silicon. EOMs based on this effect generally use $LiNbO_3$ crystals, which have yet to be integrated with CMOS manufacturing [92]. As for Kerr effect, while it is present in silicon, it is too weak for optical modulation [93].

Thermo-optical effect

Variations of silicon waveguides effective refractive index can be obtained through local changes in temperature. Since the local temperature of silicon can be controlled by Joule effect, through current injection in a nearby resistive layer, a modulation structure can easily be implemented. Those variations results from three major effects: changes in the distribution functions of carriers and phonons, temperature-induced shrinkage of the bandgap, and thermal expansion/contraction of the atomic lattice. The absorption coefficient will also be affected, but the variations are considered low as compared to those of the refractive index [92]. Therefore the thermo-optical effect is used in resonant or interferometric structures. The effect of temperature on silicon refractive index is evaluated by its thermo-optic coefficient. At room-temperature and in the $1.3\mu m$ wavelength region its value is given by [94]:

$$\frac{dn}{dT} = 1.94 * 10^{-4} K^{-1} \quad (II.12)$$

While this effect is the most efficient in silicon, it is unsuitable for high-speed modulation, due to large transition times from ON to OFF state (typically some microseconds) [95], while our targeted bit rates are in the

Gb/s range (with $T_{bit} < 1\text{ns}$). Nevertheless, this effect can be useful to set the phase difference of an interferometric structure (as discussed in **section II-2.3**) or shift the peak wavelength of a resonator.

Franz-Keldysh effect and quantum-confined Stark effect

By applying an electrical field on a material, it is possible to modify the absorption of a material, and thus create EAMs. Once applied, the electric field will distort the energy bands, resulting in electrons tunnelling across the bandgap. Electrons in the valence band will absorb photons with a lower energy than the bandgap energy, resulting in a shift of the absorption spectrum toward higher wavelengths [5]. This phenomenon is known as the Franz-Keldysh effect (FKE). FKE also involves a variation of refractive index, but this variation is low compared to the one on the absorption coefficient. Moreover, FKE is more effective for materials with a direct bandgap, and at wavelengths close to the absorption edge [92]. A similar effect exists in multiple quantum-well (MQW) structures, and is referred as the quantum-confined Stark effect (QCSE). Using these structures, the variation of absorption coefficient can be enhanced, and the wavelength of operation can also be controlled. While QCSE devices are more popular, their effect is even more confined to the proximity of the absorption edge – reducing their optical bandwidth –, and they are also more sensitive to the temperature as compared to FKE devices [3].

Silicon being an indirect bandgap material, FKE or QCSE cannot be used directly for optical modulation. However, as for the lasers, III-V materials can be used to form EAMs on silicon. Inspired by the devices made on bulk III-V, several successful demonstrations of hybrid III-V on silicon EAMs have been realized in the $1.3\mu\text{m}$ [96] or $1.55\mu\text{m}$ [97] wavelength regions, both with an ER close to 10dB at a 50Gb/s bit rate, for a $2V_{pp}$ (peak-to-peak) voltage swing, and device IL around 5dB. Schematic views of these EAMs are displayed on **Figure II-3**. This approach is all the more appropriate that to form a high-speed transmitter with a hybrid III-V on silicon laser, III-V stacks will be bonded on top of the silicon waveguides. The simplest case would then be to use the same stack for both lasers and EAMs.

However, this approach is not trivial. First of all, the excellent results presented in [96], [97] used InGaAlAs-based MQW, instead of InGaAsP that we will use for the hybrid III-V on silicon lasers, as discussed in **section III-1.1**. Indeed, due to its larger conduction band offset, InGaAlAs produces a stronger QCSE, and thus a higher extinction ratio than InGaAsP [98]. Secondly, the optimal number of wells used for the EAMs obeys to different rules from those used to design the lasers. Indeed, a thin intrinsic layer (thus a few number of wells) gives a higher extinction ratio, but leads to large capacitance, which strongly limits the speed of the device [96]. Therefore, optimized EAMs generally have more than 10 wells in their structure [96]–[99]. Thirdly, the targeted bandgaps are not the same for EAMs and lasers. Indeed, for a fixed laser wavelength, the initial absorption spectrum must be shifted at a lower wavelength to be close to the absorption edge and maximize the efficiency. Therefore the PL wavelength of the EAMs is generally shifted at 1247nm for $1.3\mu\text{m}$ operation [96], and 1478nm for $1.55\mu\text{m}$ operation [97], [98]. Finally, to reduce the capacitance of the devices, the wells width must be reduced with an under-etching step [99], which is difficult to control efficiently [98]. The devices must also be cladded in SU8 to reduce the parasitic capacitances and fill the voids left by the under-etching.

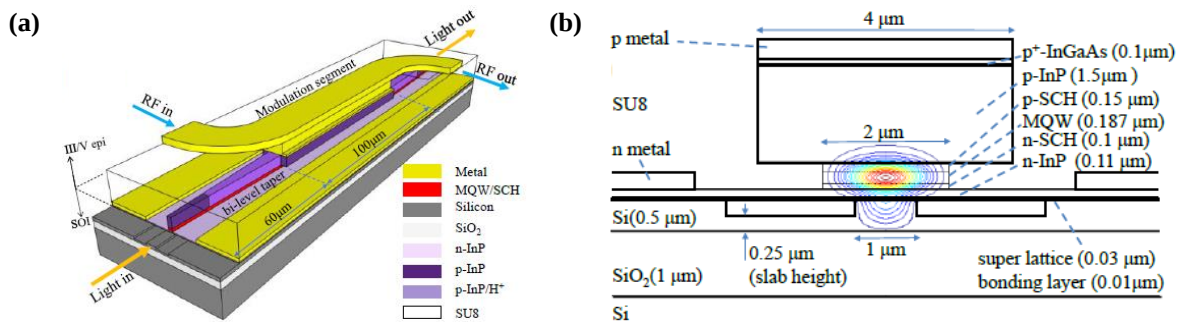


Figure II-3: Schematic (a) Top and (b) cross-sectional views of a hybrid III-V/Si EAM. Images from [97].

Therefore, to use the same III-V stack for the lasers and the EAMs, a trade-off will be necessary, since both components cannot be optimized separately. A possible solution for the bandgap issue is to use quantum-well intermixing (QWI) to create several bandgaps with the same III-V stack. Transmitters have been realized using

this method, but with limited efficiency due to the difficulty to adjust the bandgaps with QWI [100]–[102]. Another possibility is to rely on multiple-die bonding, and to use different III-V stacks adapted to the device [73]. However the dies have to be relatively close to each other and carefully aligned, which increases the fabrication process complexity.

An alternative solution to form EAMs on silicon is to use Ge, which is close to a direct bandgap material (as explained in **section I-2.3**), and is also compatible with CMOS fabrication processing. Such device based on FKE have recently shown great performances with an ER up to $3.3dB$ at a $56Gb/s$ bit rate, for a $2V_{pp}$ voltage swing, and device IL around $5.5dB$ (see **Figure II-4(a)**) [103]. However, it is limited to operations around $1610nm$. Studies are still ongoing to obtain a Ge/SiGe QCSE modulator operating in the $1.3\mu m$ wavelength region (as the one displayed on **Figure II-4(b)**), but are still lacking in term of efficiency [104], and high-speed modulation has not been demonstrated so far [105].

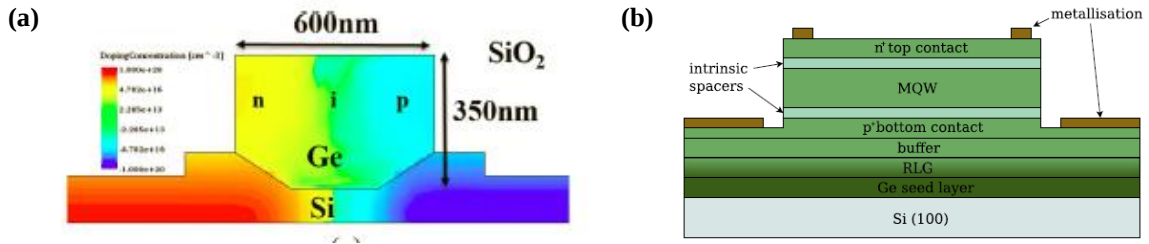


Figure II-4: Schematic cross-sectional views of (a) a Ge-based FKE modulator [103] and (b) a Ge/SiGe QCSE modulator [104].

Free-carrier plasma dispersion effect

While silicon does not have a strong electro-optical or electro-absorption effect in the $1.3\mu m$ wavelength region, it is possible to modify its refractive index and absorption coefficient by introducing free carriers (electrons and holes) coming from impurity ionization, using ion implantation. By controlling the spatial distribution of free carriers in a silicon waveguide, one will induce variations of effective refractive index and absorption coefficient, and a modulator structure can be formed. This effect is usually known as the free-carrier plasma dispersion effect (FCPDE). The changes on refractive index and absorption coefficient can be theoretically described using a Drude model, but this does not perfectly fit with the measurements [105]. Thus, studies tend to rely on semi-empirical relations for silicon introduced by Soref, which – in the $1.3\mu m$ wavelength region – are [92]:

$$n(x, y, z) = n_{Si} - [6.2 * 10^{-22} \Delta N_e(x, y, z) + 6 * 10^{-18} \Delta N_h^{0.8}(x, y, z)] \quad (II.13)$$

$$\alpha(x, y, z) = \alpha_{Si} + [6 * 10^{-18} \Delta N_e(x, y, z) + 4 * 10^{-18} \Delta N_h(x, y, z)] \quad (II.14)$$

Where n_{Si} and α_{Si} are respectively the refractive index and absorption coefficient of undoped silicon in the $1.3\mu m$ wavelength region. ΔN_e and ΔN_h are respectively the variations of electrons and holes concentrations (in cm^{-3}). These relations have been re-evaluated using more recent experimental data, but are not radically different from their previous versions [106]. Therefore, **Eqs. (II.13)-(II.14)** are considered in this study. While FCPDE enables variations of phase and absorption, modulators based on this effect are generally EOMs. However, the absorption induced by free carriers is not negligible, and add optical losses which must be taken into account.

Three mechanisms are usually considered to control the spatial distribution of free carriers via an electrical signal:

1. **Carrier injection:** Structures based on carrier injection use a $p-i-n$ diode, created through ion implantation in a rib silicon waveguide. When the device is forward biased, free carriers are injected into the waveguide, and change the phase of the light propagating through the device. An example is shown on **Figure II-5(a)**. Amongst the three mechanisms, carrier injection is the most efficient [107], with low-frequency $V_{\pi}L_{\pi}$ below $0.03V.cm$ [108], [109]. However, the E/O bandwidth of devices based on carrier injection is limited by minority carrier lifetime, and is typically in the hundreds of megahertz [105]. This issue is generally solved by employing pre-emphasis driving signals, which involves a distortion of the driving signal, to compensate for the large switching times [108]–[110]. Using this method, modulators based on carrier injection can reach

bit rates of 50Gb/s [108], [109]. Nevertheless, the use of pre-emphasis is not suitable, since it increases the complexity of the driving circuit. However, it can be noted that a carrier injection modulator has been recently demonstrated with an E/O bandwidth of 17GHz , thanks to a passive equalizer circuit, which enabled a 25Gb/s bit rate without pre-emphasis [111].

2. Carrier accumulation: Structures based on carrier accumulation use a thin insulating layer positioned in the waveguide between two semiconductor layers. When the device is biased, free carriers accumulate on both sides of the barrier, and change the phase of the light propagating through the device. An example is shown on **Figure II-5(b)**. $V_{\pi}L_{\pi}$ below $0.2\text{V}\cdot\text{cm}$ have been reported for carrier accumulation devices [112]. Unlike carrier injection devices, the speed of carrier accumulation devices is not limited by an intrinsic physical effect, and while having a large capacitance due to their thin oxide layer, they are able to reach higher E/O bandwidths [105]. Such devices have been demonstrated in the $1.3\mu\text{m}$ wavelength region at a 40Gb/s bit rate, with an ER of 8dB , for a $1V_{pp}$ voltage swing [112]. However, carrier accumulation devices are also more complex to fabricate. Indeed, they generally rely on the use of deposited polysilicon, which also adds a lot of optical losses, has explained later in **section IV-1.2**. The thickness of the oxide layer is also critical to ensure a good trade-off between efficiency and speed [113]. Therefore, there has been very few demonstrations so far of carrier accumulation devices able to realize high-speed modulation above 10Gb/s [112], [114], [115].

3. Carrier depletion: Structures based on carrier depletion use a p - n diode, formed by ion implantation in a rib silicon waveguide. When the device is reversely-biased, the depletion region widens, removing free carriers, and changing the phase of the light propagating through the device. An example is shown on **Figure II-5(c)**. Carrier depletion devices are the least efficient of the three mechanisms, with $V_{\pi}L_{\pi}$ usually above $1\text{V}\cdot\text{cm}$ [116]–[128]. However, they have also less limitations in terms of speed, since their intrinsic time response is in the picoseconds range [129], [130], and their capacitance is lower than for the carrier accumulation devices [107]. Thus several devices based on carrier depletion have been demonstrated several with bit rates ranging from 40 to 60Gb/s [117]–[128]. The main weakness of these devices is their low efficiency, which requires millimeter-long modulators if an interferometric structure is used.

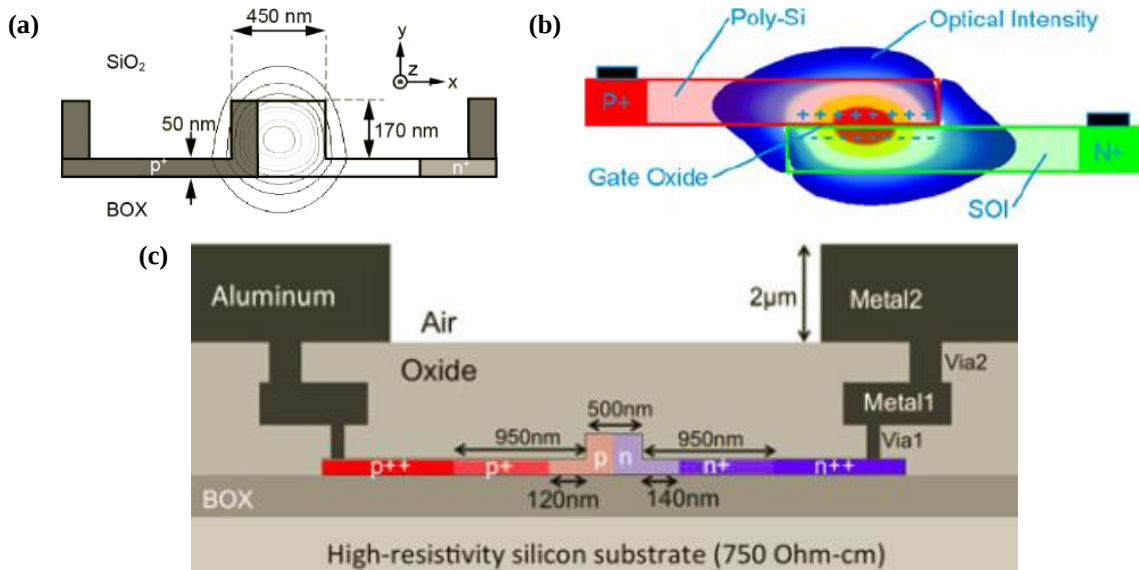


Figure II-5: Schematic cross-sectional views of modulators based on (a) carrier injection [110], (b) carrier accumulation [112], and (c) carrier depletion [127].

Conclusion and choice

In this section, the different physical effects available in silicon were presented with their pros and cons. Among them, Pockels and Kerr effects are not suited for modulation in the $1.3\mu\text{m}$ in the wavelength region. The thermo-optical effect is efficient but more suited for low-speed phase shifting.

FKE and QCSE in hybrid III-V on Si structures might be seen as more appropriate approaches, since III-V stacks will be bonded on silicon for the lasers of the integrated transmitter. However, realizing lasers and EAMs with the same III-V stack would require a trade-off, since both structures cannot be optimized on the same time. Even if solutions such as multiple die bonding exist, they are not straight-forward, and require careful alignment. Moreover, using III-V material for both devices would require more developments on the III-V processing, while – as it was explained in **section I-2.1** – one of the main interests of silicon photonics is to transfer the maximum of functions in the silicon, to benefit from its excellent fabrication controls. Therefore, hybrid III-V on silicon EAMs were not chosen for the transmitter. Ge-based EAMs have recently shown excellent results, but they are not available in the $1.3\mu\text{m}$ in the wavelength region.

Thus, the only choice left is FCPDE. While carrier injection devices are efficient, they also require pre-emphasis on the electrical signal to work at bit rates above 1Gb/s . Carrier accumulation devices can be fast and efficient, but their fabrication is more complex, and few demonstrations have been realized up to now. Therefore, since the goal is to integrate the modulator with a laser, modulators based on the more common carrier depletion were chosen, not for its efficiency, but for their higher speed and fabrication simplicity.

II-1.4. Device structures used for carrier depletion devices

Once the carrier depletion mechanism has been chosen for the modulator, the next step is to choose the structure for the modulator. Since FCPDE is used to form EOMs, a resonant or interferometric structure is needed to obtain a modulation of intensity. Two types of structures can generally be found in the literature and are presented here: ring resonators and Mach-Zehnder interferometers.

Ring resonators

Resonant structures are generally used for filters, lasers or sensors, but are also suited for modulation. Amongst the resonant structures, the ring modulator (RM) is the most commonly used. Generally, it consists in a silicon ring waveguide associated with a straight waveguide. The diode used for phase modulation is embedded in the ring waveguide. An example of RM is displayed on **Figure II-6(a)**. Light enters through the straight waveguide, and is partially coupled to the ring. If after one round-trip in the ring the phase of the coupled light is an integer multiple of 2π , it will resonate, and result in a transmission spectrum with sharp dips corresponding to optical resonances as shown on **Figure II-6(b)**. Therefore, by changing the phase through carrier depletion, the resonant wavelength is shifted and amplitude modulation can be obtained (see **Figure II-6(c)**). RMs exhibit the advantages of compact footprint, require low operational voltages to offer a large extinction ratio, and are able to reach very high data rates. Recently, a 56Gb/s bit rate was reached, using a silicon RM based on carrier depletion, with an ER of 4dB , for a $2.5V_{pp}$ voltage swing [131].

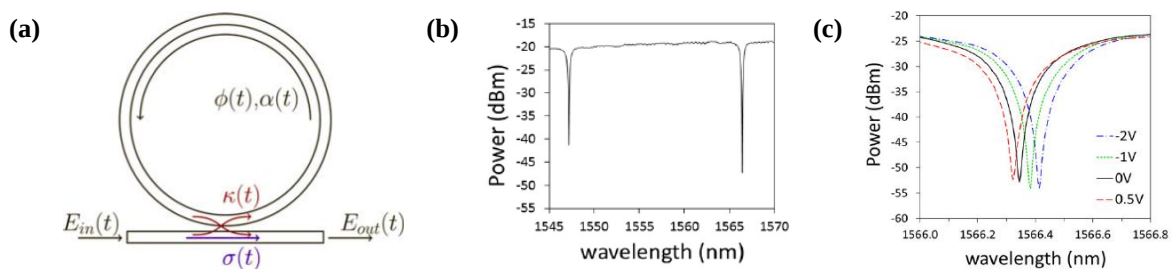


Figure II-6: (a) Top schematic view [132], and (b)-(c) transmission spectra [131] of a ring modulator.

Being resonant structures, RMs have a reduced optical bandwidth and are thus extremely sensitive towards temperature and fabrication variations [133]. Thus, they always require static phase shifters (generally obtained by thermo-optic effect) with a tight control to set them at the correct operating wavelength.

Mach-Zehnder interferometers

The Mach-Zehnder modulator (MZM) is the most commonly used structure for carrier depletion devices. It is composed of an optical splitter, two straight waveguides (also known as arms), and an optical combiner. An example of silicon MZM is shown on **Figure II-7(a)**. The diode used for phase modulation is generally

embedded in both arms, as depicted on **Figure II-7(b)**. Light passes through the optical splitter, propagates in the two arms, and is then recombined in a single waveguide. According to the phase difference between the two arms, the signals will interfere differently, leading to a change in the intensity of the output signal.

If the arms have the same length, the MZM is referred as symmetric, and else as asymmetric. If the waveguides are identical, the splitter/combiner perfectly splits/combines the light in half, and only the losses coming from the free carriers are considered, then the field magnitude out of the MZM can be expressed as:

$$E_{MZM} = \frac{E_{IN}}{2} e^{-\frac{\tilde{\alpha}(V_{bias1})}{2} L_m} e^{-j \frac{2\pi \tilde{n}(V_{bias1})}{\lambda} L_m} + \frac{E_{IN}}{2} e^{-\frac{\tilde{\alpha}(V_{bias2})}{2} L_m} e^{-j \frac{2\pi \tilde{n}(V_{bias2})}{\lambda} L_m} e^{-j \frac{2\pi \tilde{n}_{Si}}{\lambda} \Delta L} \quad (II.15)$$

Where E_{IN} is the magnitude of the input field, L_m is the length of the MZM active region, ΔL is the length difference between the arms and $\tilde{n}_{Si}$ is the effective index of the undoped silicon waveguide. The transmission of the MZM can then be expressed by raising to the square the modulus of electric field, and normalize it by the input intensity. If a bias is applied on one of the arm, and if the variations $\Delta \tilde{\alpha}$ are neglected, the transmission becomes:

$$T_{MZM} = e^{-\tilde{\alpha}(0)L_{act}} \cos^2 \left(\Delta \varphi(V_{bias}) + \frac{\pi}{\lambda} \tilde{n}_{Si} \Delta L \right) \quad (II.16)$$

In the case of an asymmetric MZM ($\Delta L \neq 0$), when there are no applied biases (which is equivalent to $\Delta \varphi = 0$), the phase difference between the two arms depends on the wavelength, resulting in an interferometric spectrum. By applying a reverse-bias on one of the arms, the phase difference between the arms is modified, resulting in a shift of the interferometric spectrum, as the one depicted on **Figure II-7(c)**. In the case of a symmetric MZM ($\Delta L = 0$), there is no interferometric spectrum, and when a reverse-bias is applied on one of the arms, the phase difference between the arms results in a modulation of amplitude over a large wavelength span, as shown on **Figure II-7(d)**.

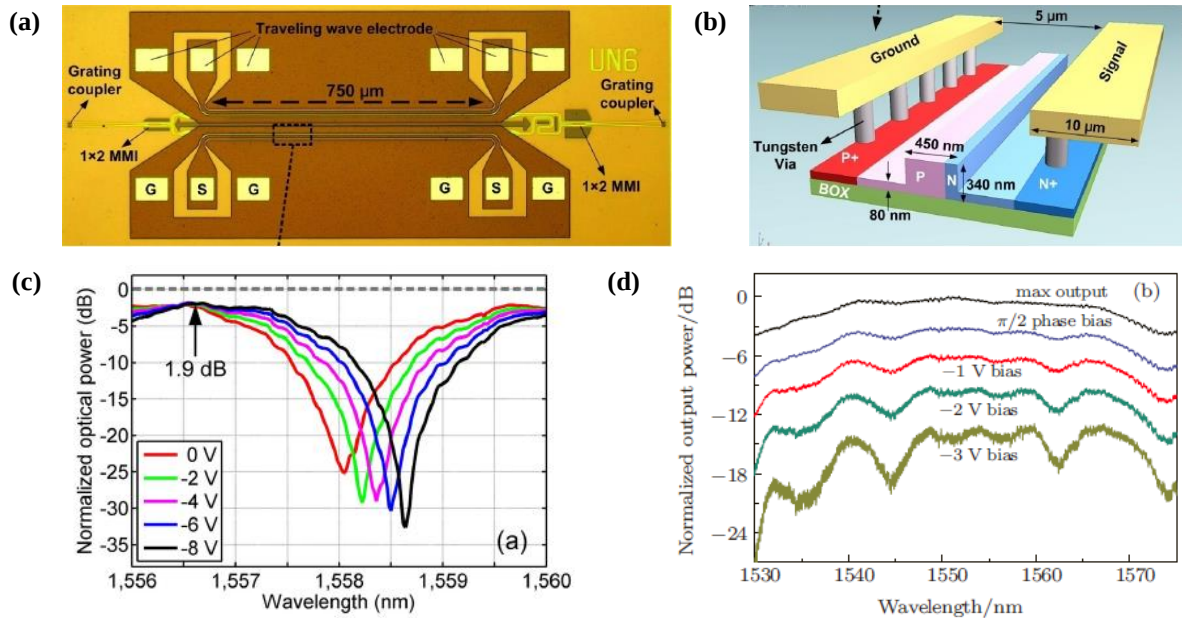


Figure II-7: (a) Top microscopic view of a silicon MZM. (b) Schematic view of the p - n junction embedded in its arms. (c) Transmission spectrum of an asymmetric MZM. (d) Transmission spectrum of a symmetric MZM. Images (a)-(c) are from [123], image (d) is from [134].

Amongst the different types of modulators, Mach-Zehnder modulators have the largest optical bandwidth, and are very tolerant to fabrication and temperature variations coming from their environment [135]. Numerous demonstrations have been realized at $1.55 \mu\text{m}$ with bit rates between 40 and 60 Gb/s [117]–[127], and more

rarely in the $1.3\mu\text{m}$ wavelength region at a 50Gb/s bit rate [128]. The main issue with MZMs based on carrier depletion is their length. Indeed, due to the relatively weak efficiency achievable via carrier depletion, MZMs generally require millimeter-long sections to reach an acceptable extinction ratio. This may result in an increase of the optical transmission loss and a decrease of the bandwidth due to a large capacitance, if the device is not carefully designed.

Finally, the MZM structure was chosen for its inherent robustness towards fabrication and temperature variations. Moreover, the modulator must be integrated with a hybrid III-V on silicon laser, which is not as convenient as a benchtop laser in terms of wavelength control, and which might be an issue for resonant type structures.

II-1.5. Target performance of the modulator

As explained in **section I-4.2**, our objective is to realize an integrated transmitter for silicon photonics, with specifications based on the 100GBASE-LR4 norm. However, most of these specifications (such as OMA) are dependent on the optical power provided by the integrated laser, which is not known at first for the future transmitter. Additionally, the co-integration of hybrid III-V on silicon lasers and silicon modulators imposes limitations on the design, due to the fabrication methods of the laser, which are not standard for silicon photonics. For these reasons, the following requirements are set for the modulators.

The silicon MZMs are based on a thin SOI layer (300nm) to be compatible with the silicon photonics fabrication platform of both CEA-LETI and STMicroelectronics [11]. In the current integration of the hybrid III-V on silicon lasers (detailed in **chapter III**), the metallic electrodes are formed by metal deposition and lift-off. Multi-level interconnections which are traditionally used in microelectronics and silicon photonics technologies are not available. Therefore, the silicon modulators used in the transmitters can only use one metal level for their electrodes. In order to fulfil the specifications described in **Table I-2**, the modulators must operate in the $1.3\mu\text{m}$ wavelength region, and be fast enough to reach a data rate of 25Gb/s . Insertion losses must be kept as low as possible to reduce the output power needed for the laser, while maximizing their extinction ratio. Finally, in order to eventually use a CMOS driving circuit to control the MZMs, their input voltage sweeps are limited to reasonable levels, such as $2.5V_{pp}$ on each arm.

II-2. Silicon Mach-Zehnder modulator design

While MZMs modulators based on carrier depletion have been demonstrated several times in the literature for high-speed modulation, each design is different, and the essential details such as the doping conditions to form the p - n junction are often not provided. Therefore, it was necessary to design our own modulators, which would correspond to our specifications. The aim of this section is to thoroughly introduce the design of the silicon MZM. At first, the design of the p - n junction is detailed, followed by the electrodes of the modulator, and finally the complete structure.

II-2.1. Junction design

The MZM comprises two parallel silicon rib waveguides, embedded in p - n junctions. The dimensions of these waveguides are based on the dimensions of passive waveguides in the CEA-LETI fabrication platform, which are 300-nm -thick and 400-nm -wide, while the slab thickness is 150nm . This way, there is no need for a tapering section between passive and doped waveguides, or an additional etching step during fabrication.

The junction is a key element of the device, since it determines its efficiency and its optical losses. Different p - n junction profiles exist and are supposed to enhance the performances of the MZM, such as p - i - p - i - n [119], interdigitated [120]–[122], or even S-shaped [136]. Finally, we decided to use a lateral p - n junction, which does not require complex processing steps, and is easier to design.

According to **Eq. (II.13)**, the change of the refractive index real part is more effective by using holes rather than electrons (for concentrations below 10^{20}cm^{-3}). Moreover, according to **Eq. (II.14)**, holes also bring less optical losses than electrons. In order to foster the effect of the holes, the junction is shifted at 100nm from the center of the waveguide. All these considerations give the structure depicted in **Figure II-8**.

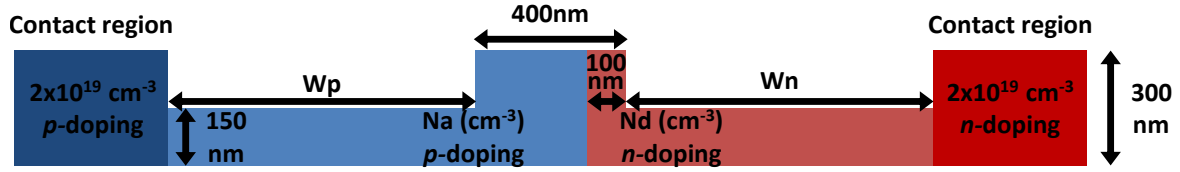


Figure II-8: Transversal schematic view of the silicon MZMs doped waveguide.

In order to electrically control the diode, contact regions are located at each side of the p - n junction. These contact regions are heavily doped in order to reduce the access resistance of the modulator, but also to reduce the contact resistance with the metal deposited on top of it. The doping concentrations of both p -doped and n -doped contact region are fixed at $2 \times 10^{19} \text{ cm}^{-3}$. While these contact regions must be located as close as possible to the junction to reduce the access resistance, they will also bring additional optical losses if too close to the fundamental mode propagating in the silicon waveguide, depicted on **Figure II-9(a)**. Therefore the parameters needed to complete the design of the junction are:

- W_p and W_n , which represent the distances between the edges of the rib waveguide and the p -doped and n -doped contact regions.
- N_a and N_d , which are the acceptor and donor doping concentrations of the p - n junction.

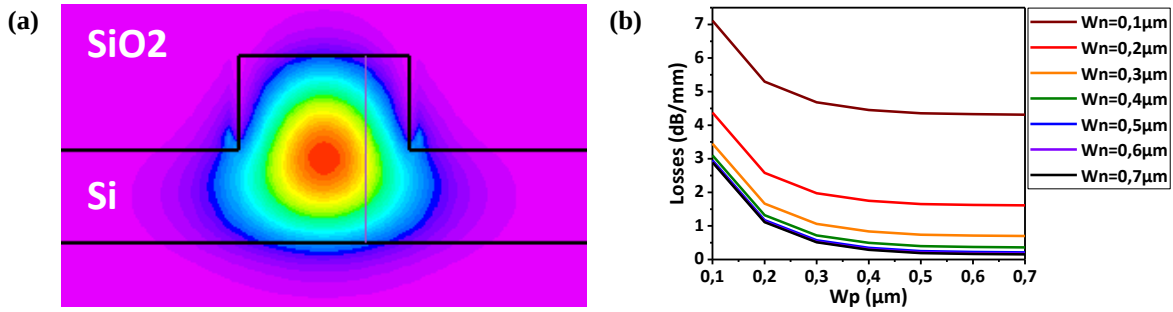


Figure II-9: (a) Fundamental TE mode propagating in the silicon waveguide. (b) Additional optical losses coming from the highly doped regions of the silicon MZMs.

To find the shortest values of W_p and W_n without adding optical losses, a model is created via the simulation software Atlas (from Silvaco). A structure similar to the one on **Figure II-8** is used without the p - n junction at the center of the waveguide. The effective refractive index and transverse modal loss can be extracted using Atlas, which uses finite-element modelling (FEM) to solve Helmholtz equation. In the simulation, only the free-carrier absorption is considered and the doping-related losses (DRL) can be expressed using (with $\tilde{\alpha}$ in cm^{-1}):

$$\text{DRL} [\text{dB/mm}] = -10 * \log_{10} \left[e^{\left(-\frac{\tilde{\alpha}}{10} \right)} \right] \quad (\text{II.17})$$

Simulations are run by sweeping the values of W_n and W_p . The result is summarized in **Figure II-9(b)**. As expected, DRL due to the contact regions are large when the contact regions are too close to the rib, but become negligible when W_n and W_p are above 500 nm . Therefore, in order to limit the risks due to fabrication variations, W_p and W_n have been fixed at 800 nm . A smaller value would reduce the access resistance, but, as shown later in **section II-2.2**, it was not necessary given the chosen doping concentrations of our junction.

The next step is to simulate different doping concentrations for N_a and N_d , in order to find the best trade-off between phase shift efficiency and optical losses. Here two models are considered to study the junction behaviour. The first one is depicted on **Figure II-10(a)**, and only uses Atlas to create an “ideal” junction, with uniform profiles both in the X and Y directions. The second one, depicted on **Figure II-10(b)**, uses the process simulation module Athena (also from Silvaco), in order to create a “realistic” junction, with simulated implantations steps, based on Monte-Carlo simulations.

In order to ensure uniformity in the vertical direction, three successive implantations, with different energies, are realized for each doped region of the p - n junction. To reduce the risk of channelling, a thin barrier

of oxide (5nm) is deposited beforehand, and the implantations are realized with a wafer tilted at 7°. Each implantation step is done in 4 rotations of 90°, to avoid an asymmetric implantation. The implantations steps are realized in the same order and follow the same annealing process that the ones used during the fabrication as described in **section II-3**.

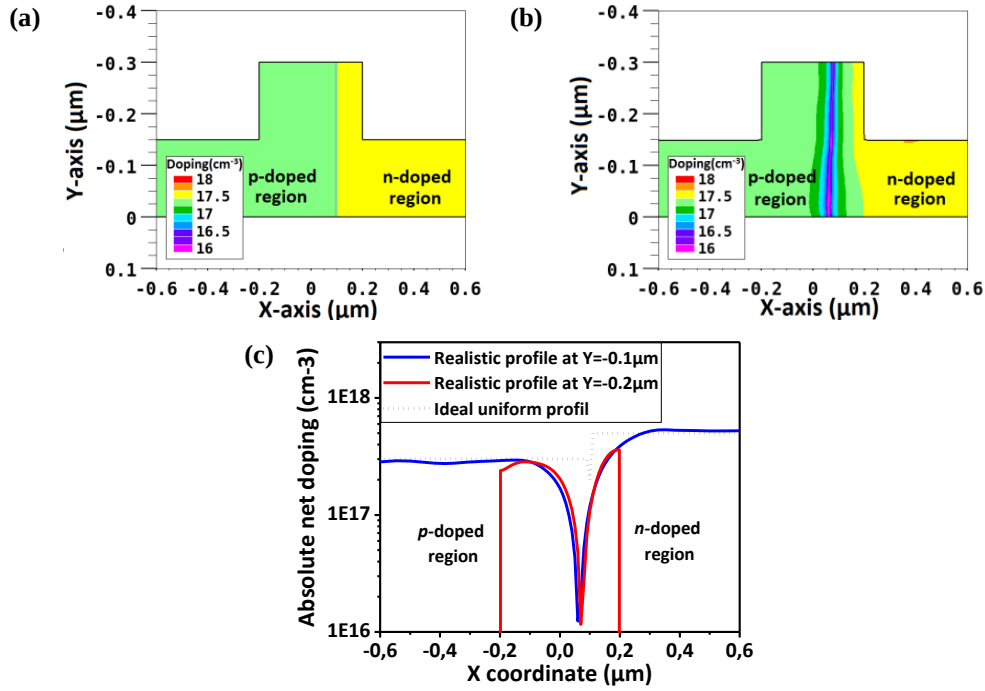


Figure II-10: Absolute net doping distribution in the waveguide, simulated with (a) an uniform doping concentration, and (b) a realistic doping profile. (c) Comparative horizontal cuts of both profiles.

A uniform profile as the one depicted on **Figure II-10(a)** is often considered for modulators, even if different from a realistic profile. The differences between the two are illustrated on **Figure II-10(c)**. While it is possible to obtain a relatively uniform doping concentration in the vertical direction with a realistic profile, the same cannot be done in the horizontal direction, especially close to the junction. Indeed the concentration of dopants drops before the limit set by the resist during implantation (located at 100nm from the center of the rib). This effect is due to the tilt of the wafer during implantation and a “shadowing” effect from the resist. Thus, with these doping conditions, it is actually difficult to obtain an abrupt junction, and an intrinsic region larger than 100nm is formed. It can also be noticed that in this example, even if the junction is supposed to be located at 100nm from the center, it is shifted towards the *p*-doped region in the realistic doping case. This effect is due to the larger concentration of *n*-dopants, and would be inverted if the *p*-doping concentration was larger.

Due to the non-uniformity of the realistic profile in the horizontal direction, the doping concentration of both *p*-doped and *n*-doped region is defined by its value outside from the junction, once it becomes constant (see **Figure II-10(c)**). The doping conditions simulated to have a doping concentration equal to $1 \times 10^{17} \text{ cm}^{-3}$ are presented in **Table I-1**. By multiplying the dose of the implantations (with the same energies) by any number *N*, it will give the doping conditions for a doping concentration of $N \times 10^{17} \text{ cm}^{-3}$.

Table II-1: Doping conditions for a uniform concentration of $1 \times 10^{17} \text{ cm}^{-3}$ in a 300-nm-thick silicon waveguide.

Doping type	Dopant	Tilt (°)	Twist (°)	Number of rotations	Dose (at/cm²)	Energy (keV)	Targeted concentration (at/cm³)
n-doping	Phosphorus	7	27	4	2.2×10^{12}	200	1×10^{17}
					1.1×10^{12}	100	
					0.55×10^{12}	35	
p-doping	Boron	7	27	4	2×10^{12}	80	1×10^{17}
					1×10^{12}	35	
					0.5×10^{12}	10	

Once the models are set, Atlas is used to simulate the device at different voltage bias. Thus, the free carrier distribution is known for any voltage bias, and the effective refractive index and transverse modal loss can be extracted. Then, using Eq. (II.4) and (II.17), it is possible to evaluate the linear phase shift and the *DRL* of the junction. A sweep is realized on several combinations of N_a and N_d for the junction. To make our choice, we take the linear phase shift at a $-2.5V$ bias, and the *DRL* at a $0V$ bias (which is a worst case, since when the free carriers are removed by the depletion, the free-carrier absorption decreases).

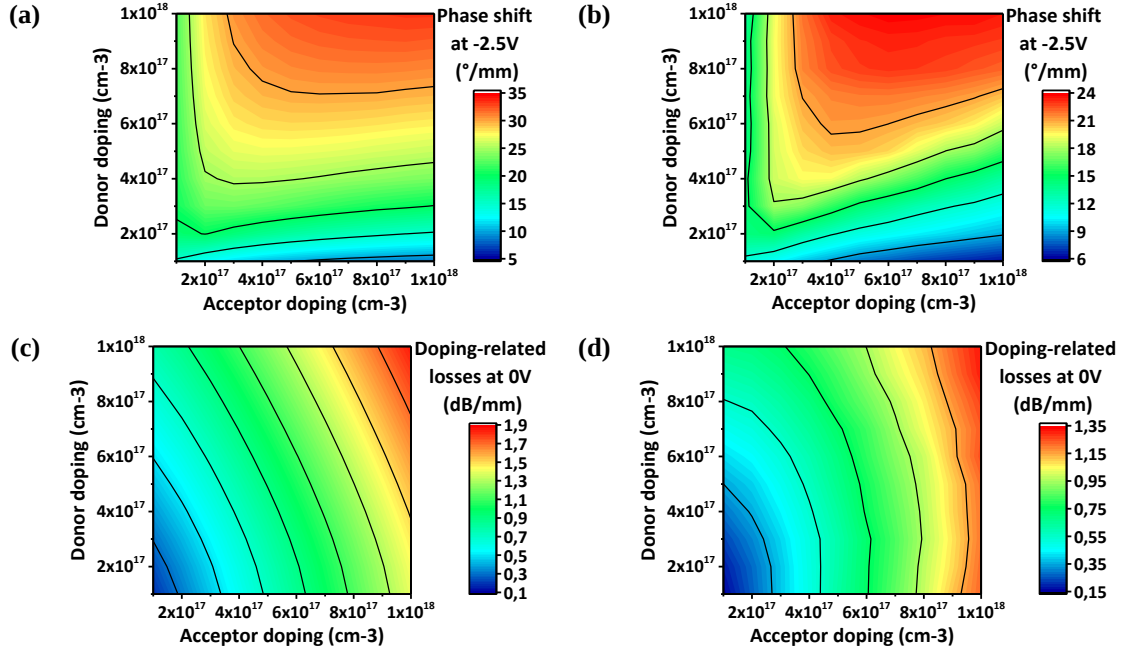


Figure II-11: Phase shift at a $-2.5V$ bias and doping-related losses coming from doping at a $0V$ bias, for (a,c) uniform doping concentrations, and (b,d) realistic doping profiles.

The results for each model are shown on **Figure II-11**. The phase shift for each combination of doping concentration is displayed on **Figure II-11(a)-(b)**, and the *DRL* on **Figure II-11(c)-(d)**. As it can be expected, both models give similar trends. In both cases, increasing the doping concentrations result in higher losses, but not necessarily in a higher efficiency. However, having a higher N_d concentration appears to have a less impact on the losses, while increasing significantly the efficiency. The difference between the two models lies in the estimated values of both phase shift and *DRL*. The model based on uniform doping concentrations is more optimistic in terms of efficiency, but also more pessimistic in terms of optical losses. This comes from the intrinsic region created during implantation, which is not taken into account in the first model. To choose the best doping concentration, and predict more precisely the performances of the junction, we rely on the realistic doping profile, which is supposed to be closer to the real device. To do so, both phase shift and *DRL* for each doping concentration combination are summarized in a single plot, shown on **Figure II-12**.

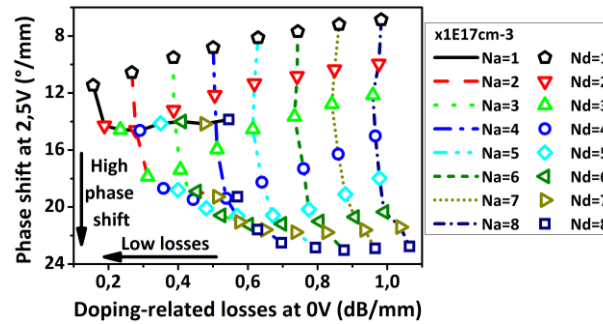


Figure II-12: Phase shift at $-2.5V$ bias versus doping related losses at $0V$ bias for realistic doping profiles. Each point represents a different combination of acceptor (N_a) and donor (N_d) atoms concentrations.

Finally, the concentrations of the p -doped and n -doped regions are respectively fixed at $4 \times 10^{17} \text{cm}^{-3}$ and $6 \times 10^{17} \text{cm}^{-3}$. This p - n junction should result in a linear phase shift $\Delta\phi \approx 21.5^\circ/\text{mm}$ at -2.5V , which is equivalent to $V_\pi L_\pi \approx 2.1 \text{V.cm}$ according to **Eq. (II.6)**, with $\text{DRL} \approx 0.63 \text{dB/mm}$ at 0V . This combination offers a good trade-off, as well as robustness towards doping variations. The expected performances of this junction and its linear junction capacitance – also extracted from Atlas, and which will be used later on – are displayed on **Figure II-13**.

Since these simulations could not be verified beforehand, and the optical losses of the other components – passive waveguides and splitter/combiner – was also unknown, we choose to implant the junction into each arm of several MZMs with different lengths for the active region, in order to find the best compromise. Three lengths were chosen: 2mm (to reduce the optical losses), 4mm and 6mm (to enhance the efficiency).

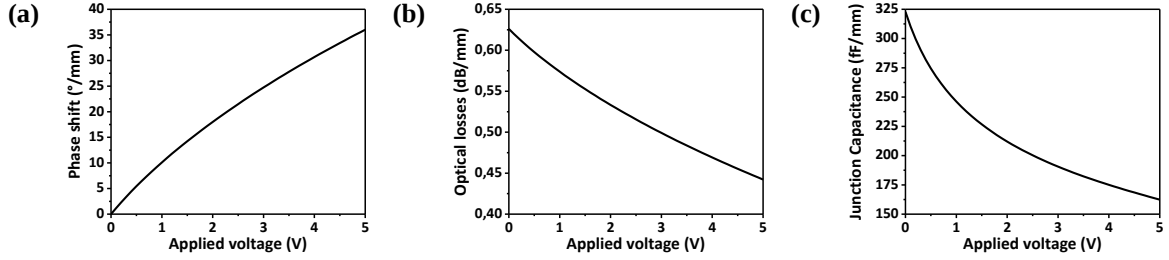


Figure II-13: (a) Linear phase shift, (b) doping-related losses, and (c) junction capacitance versus reverse-voltage bias applied on the chosen junction for the modulator.

II-2.2. Speed limitations and travelling-wave electrodes design

As we explained in **section II-1.3**, the speed of a modulator can be limited by the phenomenon used for modulation. The intrinsic response of carrier depletion devices depends on the time required for carriers to be swept out from and to return to the junction. Under a large electrical field – as it is the case in a p - n junction –, the drift velocity (v_d) reaches its saturation at 10^7cm/s . For the chosen doping concentrations, the depletion region w_d is at most 220nm wide bias at a 2.5V reverse bias according to our simulations. Therefore, the drift time is estimated at $t_d = w_d/v_d = 2.2\text{ps}$, which is more than enough for a bit rate of 25Gb/s ($T_{bit} = 40\text{ps}$).

Thus, the main factor limiting the speed of the MZMs is their capacitances. However, while the active region of a MZM can be seen as a concentrated capacitor, since its length is over 2mm , a concentrated (or lumped) design is not suitable for high-speed modulation. As a first-order approximation, we do not take into account the access resistance and other parasitic components, and only consider the diode as its junction capacitance (C_j). If the modulator is driven by a generator associated with an impedance Z_G – which is generally a resistance with a standard value of 50Ω –, the RC cut-off frequency (f_{RC}) is defined as [90]:

$$f_{RC} = \frac{1}{2\pi Z_G C_j} \quad (\text{II.18})$$

Since the targeted voltage sweep is $2.5V_{pp}$, we can fix the working bias at a -1.25V , so that the modulator will operate between 0 and -2.5V . At this working bias, the junction capacitance is approximately 235fF/mm (according to **Figure II-13(c)**). Thus, the RC cut-off frequencies are respectively 6.8 , 3.4 , and 2.3GHz for the 2 , 4 , and 6-mm -long active regions. These values are too low for the targeted 25Gb/s bit rate, especially since parasitic components have not been taken into account, and will further degrade the bandwidth.

In order to overcome the RC limitation, a travelling-wave modulator structure is generally preferred [90], [130]. In this case, the modulator is seen as a transmission line, where the capacitance is distributed over the entire device length. A transmission line model is shown on **Figure II-14(a)**. The RF signal is sent from the generator, propagates on the coplanar electrodes, and travels together with the optical signal. The transmission line is terminated by an RF load (Z_L). During its propagation, the RF wave experiences losses and reflections at the end of the line. The transmission line is defined by two characteristic parameters: the characteristic impedance (Z_c) and the complex propagation constant (γ). Those two parameters can be calculated by using an equivalent electrical model of the transmission line (displayed on **Figure II-14(b)**) and the following formulas [90]:

$$Z_c[\Omega] = \sqrt{\frac{(R + jL\omega_e)}{(G + jC\omega_e)}} \quad (\text{II.19})$$

$$\gamma [m^{-1}] = \tilde{\alpha}_{RF} + j \frac{2\pi f_e \tilde{n}_{RF}}{c_0} = \sqrt{(R + jL\omega_e)(G + jC\omega_e)} \quad (\text{II.20})$$

$$RF \text{ Losses [dB/mm]} = 20 \log_{10}(e^{\tilde{\alpha}_{RF}})/1000 \quad (\text{II.21})$$

R , L , G , and C are respectively the linear resistance, inductance, conductance and capacitance of the line. $\tilde{n}_{RF}$ is the effective index of the RF wave, and $\tilde{\alpha}_{RF}$ is the line attenuation. c_0 is the light velocity in vacuum. We also define the RF losses, which are a convenient parameter that will be used later in this section. All these parameters are frequency dependent, and change with the geometry of the travelling-wave electrodes (TWE), which must be carefully designed to maximize the bandwidth.

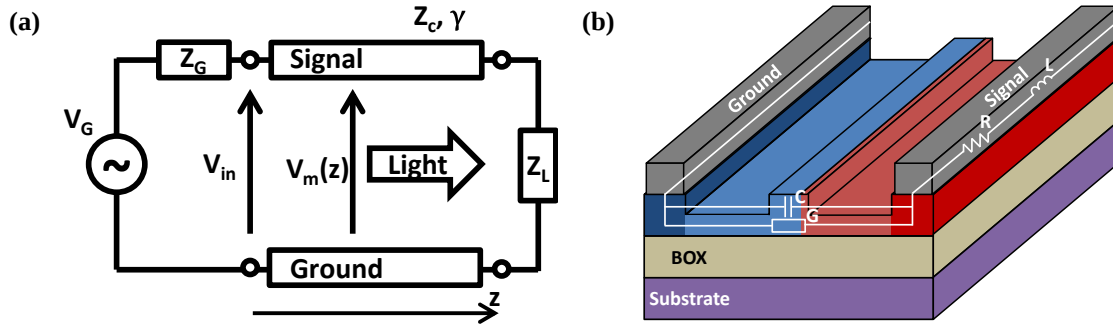


Figure II-14: (a) Transmission line model. (b) RLGC model of the modulator active region.

The frequency response of a modulator based on TWE is dominated by three main factors [90], [130]:

- 1) The velocity mismatch between the optical and RF propagating waves. It can be evaluated by comparing the group index of the optical wave (estimated at ≈ 4 given the geometry of the silicon waveguide) and the effective index of the RF wave.
- 2) The impedance mismatch between the modulator, generator and load (which are both fixed at 50Ω in our case). If too important, it will result in reflections at the input and output of the line.
- 3) The RF losses on the modulator line, which must be reduced as much as possible. If too large, the RF wave will not reach the end of the line, which will not be modulated.

By optimizing the trade-off between these three parameters, one can significantly improve the modulator bandwidth.

As explained previously in **section II-1.5**, only one metal level is allowed to define the electrodes, which limits the design options. The configuration chosen for the electrodes is a coplanar-strip signal-ground-signal (SGS) configuration, where the p - n junctions in each waveguide are mirrored, as shown in **Figure II-15**. We did not choose the more commonly used GSGSG configuration because it is asymmetric, since a p - n junction is not located between all electrodes. This asymmetry will excite undesired RF modes, and create notches in the frequency response. This issue can be overcome by using metallic connections from another metal level between the ground lines [116], [117], [127], which is not possible here. Another solution would have been to use wire bonding between the ground lines [120]. Finally, the best solution was to use the SGS configuration, which also permits to control separately each arm of the MZM in order to enhance its performances, as it will be discussed in **section II-4.4**.

We used two-dimensional FEM electromagnetic field simulations to determine the R , L , G , and C parameters for several electrodes geometries, at different frequencies. The previously designed p - n junction doping concentrations and depletion region (for several voltages) were included, with a 800-nm -thick BOX layer, a silicon substrate and a 650-nm -thick metal (AlCu) layer. After several simulations we found a good trade-off with the dimensions shown on **Figure II-15**. The characteristics of these electrodes are detailed hereafter.

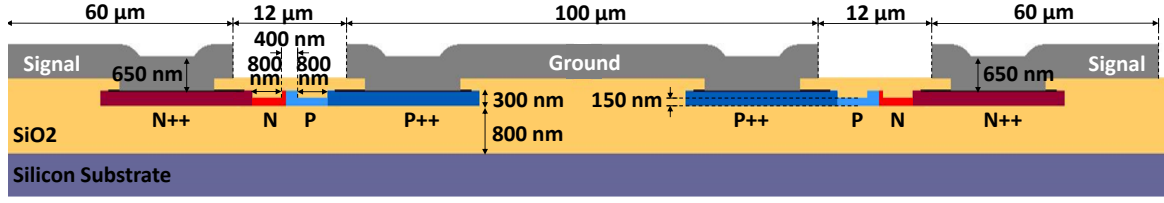


Figure II-15: Schematic transversal view of the silicon MZM, with its definitive dimensions.

RLGC parameters

The R , L , G , and C parameters of the line are displayed on **Figure II-16** for 0V and 1.25V reverse-bias voltage. The influence of the substrate resistivity is also studied, with standard resistivity (SR), and high resistivity (HR) substrates. In the simulations, the substrate resistivity is respectively $10\Omega \cdot cm$ and $750\Omega \cdot cm$.

At low frequency, the resistance of the line can be approximated by Ohm's law. When the frequency increases, the cross-section of the metal conducting the signal decreases due to the skin effect, which in turn increases the resistance (see **Figure II-16(a)**). The substrate also has an influence on the resistance, which is due to a longitudinal current inside the substrate and the silicon waveguide layer, and act as a parallel resistance [130]. Thus, by increasing the substrate resistivity, its influence on the resistance becomes negligible.

The inductance can be split in two contributions: an external inductance (related to the magnetic energy stored in the dielectric surrounding the line) and an internal inductance (related to the magnetic energy stored within the conductors). Due to the skin effect, the internal inductance decreases when the frequency increases, while the external inductance is frequency independent [90], [130]. Therefore, the inductance converges to a constant value (which depends on the electrodes geometry) when the frequency increases (see **Figure II-16(b)**).

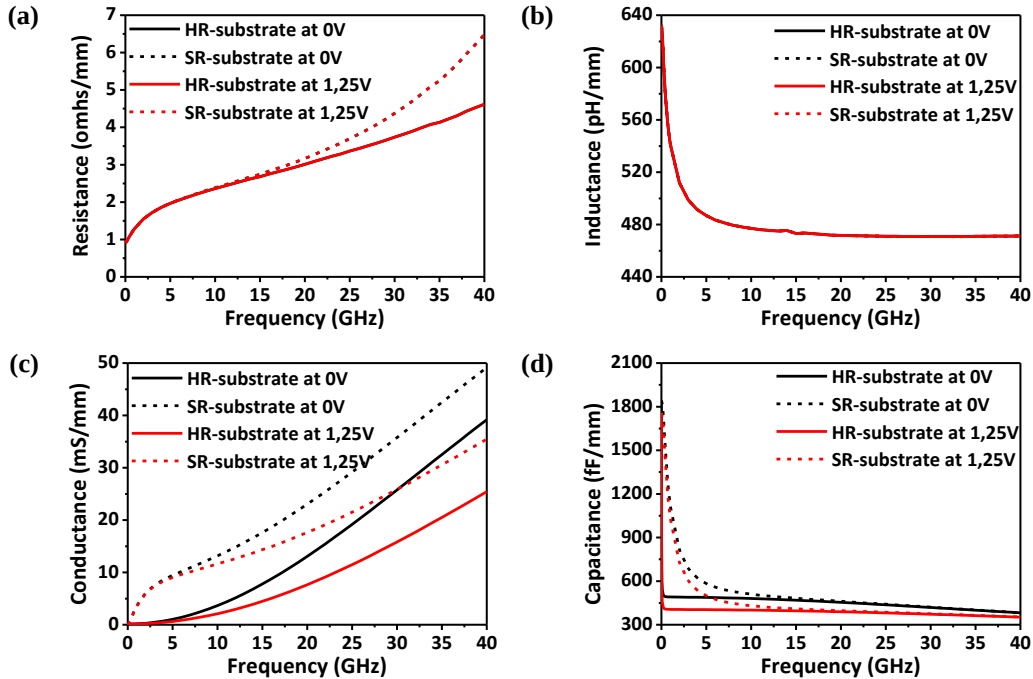


Figure II-16: Simulations of the linear (a) resistance and (b) inductance of the electrodes, and of the linear (c) conductance and (d) capacitance between the electrodes. Lines: case for a HR substrate ($750\Omega \cdot cm$). Dashed lines: case for a SR substrate ($10\Omega \cdot cm$).

In order to understand the behaviour of the conductance and capacitance of the line, it is necessary to use another equivalent electrical circuit, as the one shown on **Figure II-17** [120], [124], [130]. C_{elec} is the capacitance between the signal and ground lines, and C_{BOX} is the capacitance of the BOX. R_n and R_p are the resistances of the p -doped and n -doped silicon regions. C_{CSUB} is the capacitance between the signal metal and the substrate. C_{SUB} and G_{SUB} are respectively the substrate capacitance and conductance. Amongst the several

parasitic components, those coming from the substrate have a large influence and are known to degrade the modulator response [137]. In order to reduce the influence of the substrate, the best solution is to use a high resistivity (HR) substrate, rather than a standard resistivity (SR) one, to efficiently cut-off at low frequency the capacitances coming from the substrate. As depicted on **Figure II-16(c)-(d)**, using a HR substrate largely reduces the total conductance and capacitance between the electrodes.

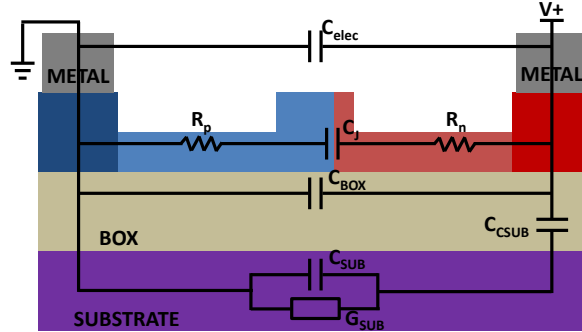


Figure II-17: Equivalent analytic circuit of the active region.

If the substrate influence can be neglected at high frequency, the model shown on **Figure II-17** can be simplified – by removing C_{SUB} , C_{BOX} , and G_{SUB} –, and the conductance and capacitance of the line can be expressed as:

$$C [F/mm] = \frac{C_J}{1 + [\omega_e(R_n + R_p)C_J]^2} + C_{elec} + C_{BOX} \quad (\text{II.22})$$

$$\frac{1}{G [S/mm]} = (R_n + R_p) \left(1 + \frac{1}{[\omega_e(R_n + R_p)C_J]^2} \right) \quad (\text{II.23})$$

Using these formulas, we can see that both the conductance and capacitance of the line are strongly dependent on the junction capacitance. Therefore, when a reverse-bias is applied, C_J decreases (see **Figure II-13(c)**), which in turn decreases both G and C , as shown on **Figure II-16(c)-(d)**. The reverse-bias does not have an effect on the resistance and inductance of the line.

To summarize, if the substrate resistivity is high enough, the resistance and inductance are only dependent on the geometry of the metal lines. The capacitance also has a dependence on the geometry of the metal lines (via C_{elec} and C_{BOX}), but is mainly dependent on the junction capacitance and on the access resistance ($R_n + R_p$). The conductance is only dependent on the junction capacitance and on the access resistance, which both depend essentially on the chosen doping concentrations for the junction.

Characteristic impedance, RF effective index and RF losses

Using **Eq. (II.19)-(II.21)**, the characteristic impedance, RF effective index and RF losses can be calculated. The results are displayed on **Figure II-18**. At a 1.25V reverse-bias, the characteristic impedance is approximately 35Ω , and the effective index is close to 4, which provides a negligible velocity mismatch.

To further understand the influence of the R , L , G , C parameters on the characteristic impedance, RF effective index and RF losses, **Eq. (II.19)-(II.21)** can be simplified. According to the values of R , L , G and C shown on **Figure II-16**, the conditions $j\omega_e C \gg G$ and $j\omega_e L \gg R$ are verified. Therefore, we are in the high-frequency regime, and the following approximations are valid [90]:

$$Z_c \approx \sqrt{\frac{L}{C}} \quad (\text{II.24})$$

$$\tilde{n}_{RF} \approx c_0 \sqrt{LC} \quad (\text{II.25})$$

$$\tilde{\alpha}_{RF} \approx \frac{R}{2Z_c} + \frac{GZ_c}{2} \quad (\text{II.26})$$

When the substrate resistivity is increased, the resistance, capacitance, and conductance of the line decrease, while the inductance remains unchanged. Using Eq. (II.24)-(II.26), it can be seen that with a HR substrate, the real part characteristic impedance increases, while the RF effective index and RF losses both decrease, as shown on Figure II-18. The use of a HR substrate is particularly essential to reduce the RF losses.

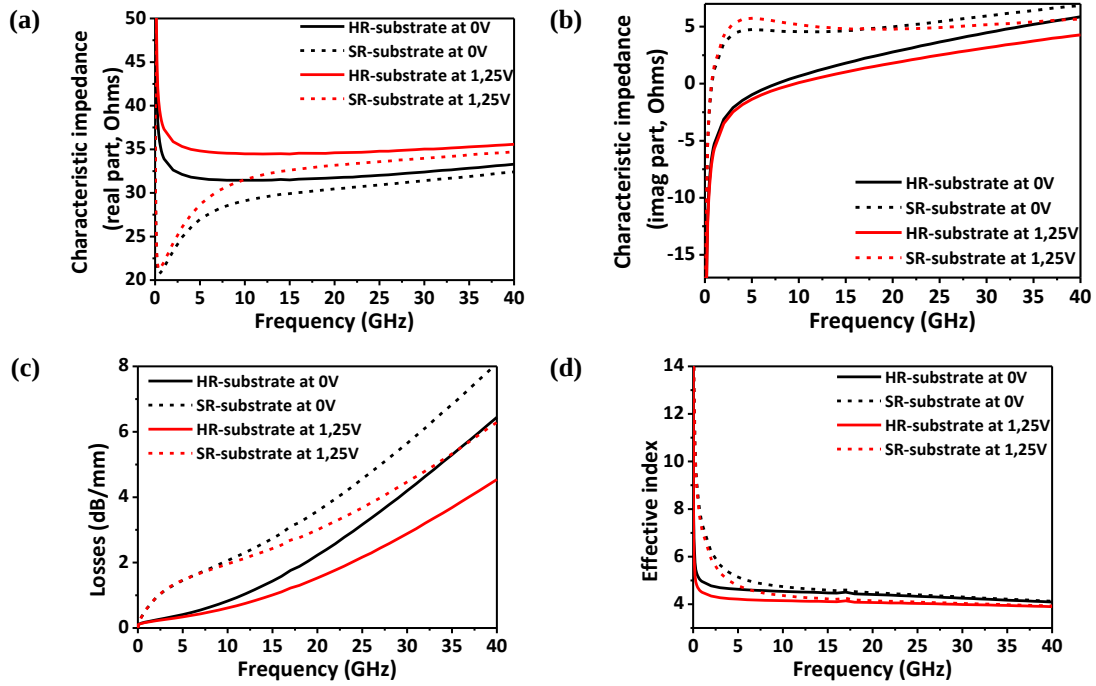


Figure II-18: Simulations of the (a) real and (b) imaginary part of the characteristic impedance, and of the (c) RF losses and (d) RF effective index. Lines: case for a HR substrate ($750\Omega \cdot \text{cm}$). Dashed lines: case for a SR substrate ($10\Omega \cdot \text{cm}$).

Using Eq. (II.24) and (II.25), it can be seen that in our configuration, it is impossible to cope with both the impedance matching (thus to increase Z_c from 35Ω to 50Ω) and the velocity matching conditions (which is already reached) at the same time. Indeed, since the capacitance of the line is mainly dependent on the doping concentrations – which were fixed before designing the electrodes –, both Z_c and $\tilde{n}_{RF}$ will mainly depend on the inductance. Since both of them are proportional to the square root of the inductance, a trade-off was mandatory.

Scattering parameters

In order to validate the trade-off, the scattering parameters (S-parameters) of the transmission line are evaluated. The transmission line is seen as a 2-port device, and is represented by a 2×2 S-matrix. The parameters S_{11} and S_{22} give the reflection of the RF signal at each port and the parameters S_{21} and S_{12} the transmission of the RF signal through the line. Since the line is symmetrical, $S_{11} = S_{22}$ and $S_{21} = S_{12}$. The S-parameters measurements are used to extract precisely the RF losses, characteristic impedance and effective index. The relations between S-parameters and line parameters can be found in [138].

The S-parameters magnitudes for each active region length are displayed on Figure II-19. The results are summarized in Table I-2. In each case, the generator and load impedance are both 50Ω . In the case of a HR substrate, and at a 1.25V reverse-bias, the S_{11} parameter is below -10dB for each length, which is an acceptable value. The reflection increases for a 0V bias or substrate resistivity, since the impedance mismatch is larger.

In the case of a HR substrate, and at a 1.25V reverse-bias, the S_{21} parameter -6dB bandwidth (which mean that the amplitude of the input signal is reduced by 50%), are above 14GHz for each active region length. It can be seen that the -6dB bandwidth for each length corresponds to the frequency where the total RF losses of the line are equal to -6dB . Therefore, the transmission of the RF signal depends mainly on the RF losses. As expected, the use of an HR substrate or a higher reverse bias drastically improves the transmission.

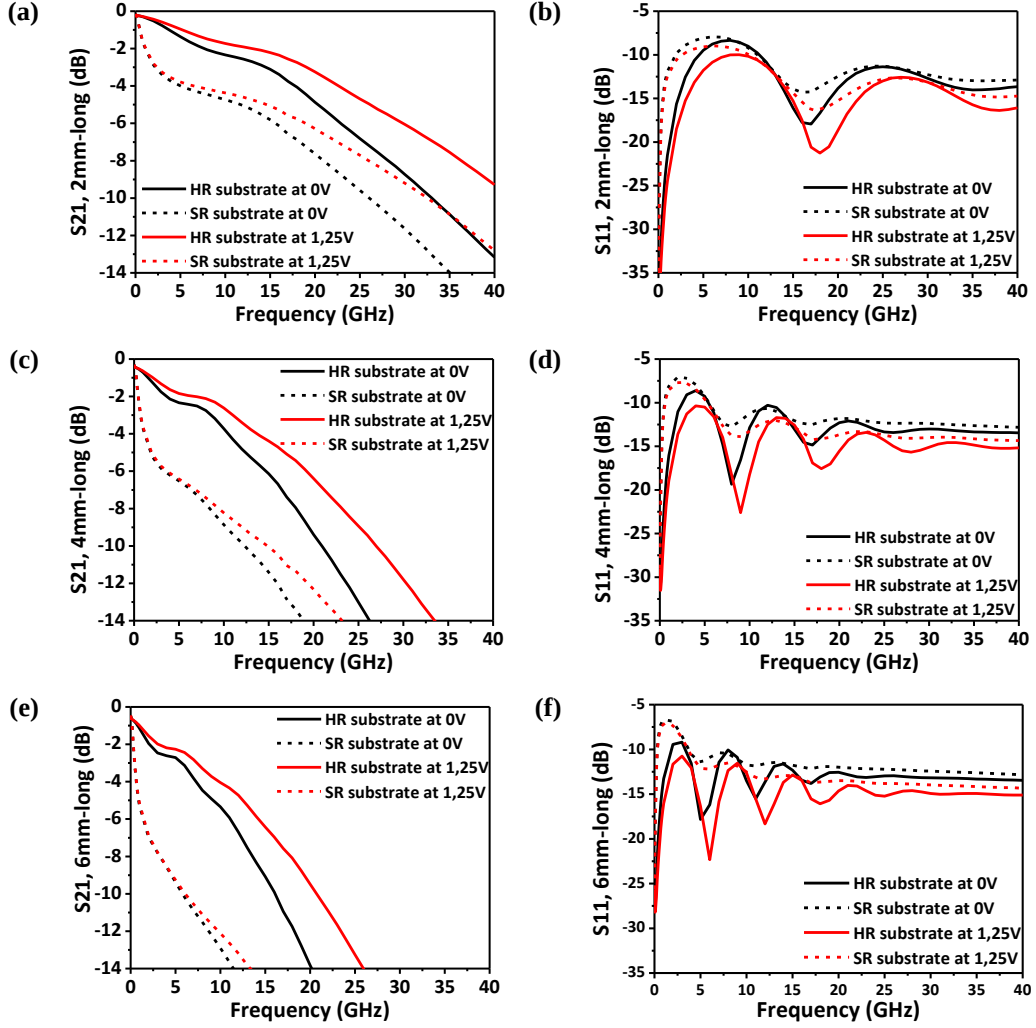


Figure II-19: Simulations of the S-parameters magnitudes for the different lengths of the active region: 2-mm-long (a) and (b), 4-mm-long (c) and (d), 6-mm-long (e) and (f). Lines: case for a HR substrate ($750\Omega.cm$). Dashed lines: case for a SR substrate ($10\Omega.cm$).

Table II-2: Simulated S-parameters performances with a HR substrate.

Active region length (mm)	S_{11} parameter maximum (dB)		S_{21} parameter $-6dB$ bandwidth (GHz)	
	0V bias	1.25V reverse-bias	0V bias	1.25V reverse-bias
2	-8.3	-10	22.9	29.8
4	-8.6	-10.4	14.7	19.2
6	-9.1	-10.8	11.1	14.2

Electro-optical bandwidth

Finally, the modulation depth of the modulator can also be evaluated, by using the following formula:

$$m(\omega_e) = \left| \frac{[1 + j\omega_\phi C_J(R_n + R_p)]V_{avg}(\omega_e)}{[1 + j\omega_e C_J(R_n + R_p)]V_{avg}(\omega_\phi)} \right| \quad (II.27)$$

Where V_{avg} is the average voltage between the signal and ground electrodes experienced by a photon as it travels along the modulator. It takes into account the RF losses, impedance, and velocity mismatch. Its complete expression can be found in [120] or [130]. According to the simulations of the p - n junction, C_J is estimated to $323fF/mm$ at $0V$, and $235fF/mm$ at $-1.25V$, with a series resistance of approximately $9\Omega.mm$. The calculated modulation depth for each length is shown on **Figure II-20**. In each case, the use of an HR substrate

or a higher reverse bias drastically improves the electro-optical bandwidth. The simulated electro-optical bandwidths with a HR substrate are summarized in **Table II-3**.

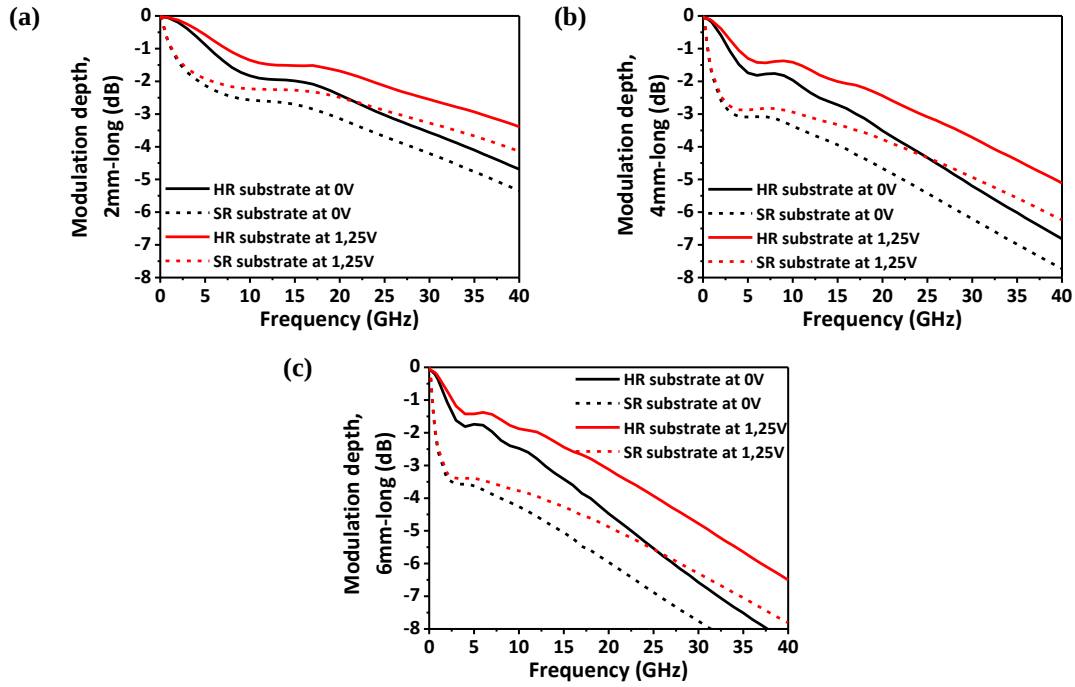


Figure II-20: Simulations of the modulation depth for the different lengths of the active region: (a) 2-*mm*-long, (b) 4-*mm*-long, and (c) 6-*mm*-long. Lines: case for a HR substrate (750 Ω .*cm*). Dashed lines: case for a SR substrate (10 Ω .*cm*).

Table II-3: Simulated electro-optical bandwidths with a HR substrate.

Active region length (mm)	Electro-optical bandwidth (GHz)	
	0V bias	1.25V reverse-bias
2	24.7	35.7
4	17.1	24.3
6	12.9	19.3

It can be noted that these values might have been improved. For instance, it might have been useful to reduce the junction capacitance, and find a better trade-off for Z_c and $\tilde{n}_{RF}$. As a first-order approximation^c, the junction capacitance can be estimated by the following equations [140]:

$$C_J(V_{bias}) = \epsilon_0 \epsilon_{Si} \frac{h_{rib}}{w_d(V_{bias})} \quad (\text{II.28})$$

$$w_d(V_{bias}) = \sqrt{\frac{2\epsilon_0 \epsilon_{Si}}{q} \left(\frac{1}{N_a} + \frac{1}{N_d} \right) \times \left[\frac{k_B T}{q} \ln \left(\frac{N_a N_d}{n_i^2} \right) - V_{bias} - 2 \frac{k_B T}{q} \right]} \quad (\text{II.29})$$

Where h_{rib} is the height of the silicon waveguide, ϵ_0 is the vacuum permittivity, ϵ_{Si} is the relative permittivity of silicon, k_B is the Boltzmann constant, T is the ambient temperature, and n_i is the intrinsic carrier density of silicon. Using these equations, it can be seen that the junction capacitance at fixed bias decreases with the acceptor and donor doping concentrations. However, the phase shift efficiency would also decrease. Additionally, reducing further the spacing between the junction and the contact regions would also reduce the access resistance, which in turn would decrease the conductance, and thus the RF losses, but at the risk of

^c Eq. (II.28) is valid for infinite parallel plates, and does not take into account fringing fields. Eq. (II.29) is valid for ideal abrupt junctions, and is a 1-D model which no longer holds at the top and bottom of the waveguide. A more precise analytical modelling can be found in [139].

increasing the optical losses. Since these values were sufficient for a 25Gb/s modulation rate, we did not try to improve further the design.

II-2.3. Device architecture

A layout view of the MZM is displayed on **Figure II-21**. Light is coupled in and out of the device by using two grating couplers, with peak wavelengths at $1.3\mu\text{m}$. Light splitting and recombining is provided by two multi-mode interferometers (MMIs). The active regions which comprise two silicon waveguides embedded in the designed p - n junction and the electrodes (which cross-section is shown on **Figure II-15**) can be 2, 4 or 6- mm -long.

A 100- μm -long length difference of passive waveguide has also been added to form an asymmetric structure. Both symmetric and asymmetric structures are not different in terms of modulation performances, but the asymmetric structure offers the advantage to easily extract the optical group index and the phase shift, as it will be shown in section II-4.2. Nevertheless, a symmetric design having a large optical bandwidth would be more suited for a more advanced device.

A 1- mm -long section of passive waveguide is also present before the active region. As it will be discussed in section II-4.4, while the aim of the active region in each MZM arm is to provide a phase shift to modify the phase difference between the two arms, it is necessary to add a static phase difference between the two arms of the MZMs to enhance their efficiency. As explained in section II-1.3, by heating a waveguide in this passive section, its effective index of refraction – thus its phase – will be increased. Therefore, for a fixed wavelength, the static phase difference between the arms can be controlled by using a thermo-optical shifter.

However, while a static phase shifting region is mandatory for a symmetric MZM, in the case of an asymmetric MZM, the phase difference between the two arms depends on the wavelength. Thus, to test our stand-alone modulator, the static phase difference can be set by modifying the emitting wavelength of the external laser, and the static phase shifting regions are not used. Nevertheless, they will be essential for the integrated transmitter.

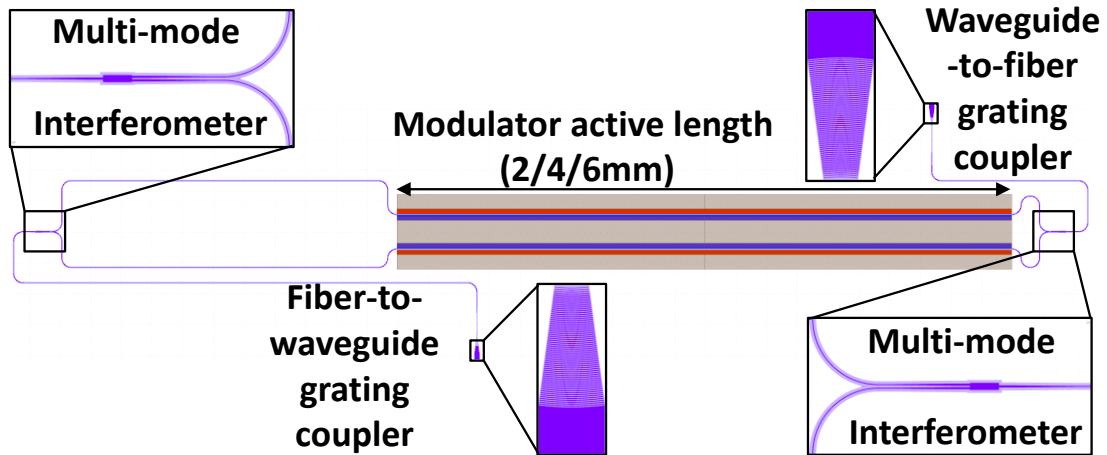


Figure II-21: Layout view of the complete MZM structure.

II-3. Silicon Mach-Zehnder modulator fabrication

The fabrication steps of the silicon MZMs are illustrated on **Figure II-22**. These devices are fabricated on 8" SOI wafers from SOITEC, with a 300-nm-thick silicon layer, a 800-nm-thick BOX, on a HR handle silicon wafer (750 $\Omega \cdot \text{cm}$). The process flow starts with the implantation steps necessary to form the *p-n* junctions. Then, the waveguides are patterned, and the contact regions are silicided in order to reduce the contacts resistance, and thus the access resistance. Finally, the devices are encapsulated and planarized, before forming the modulators electrodes.

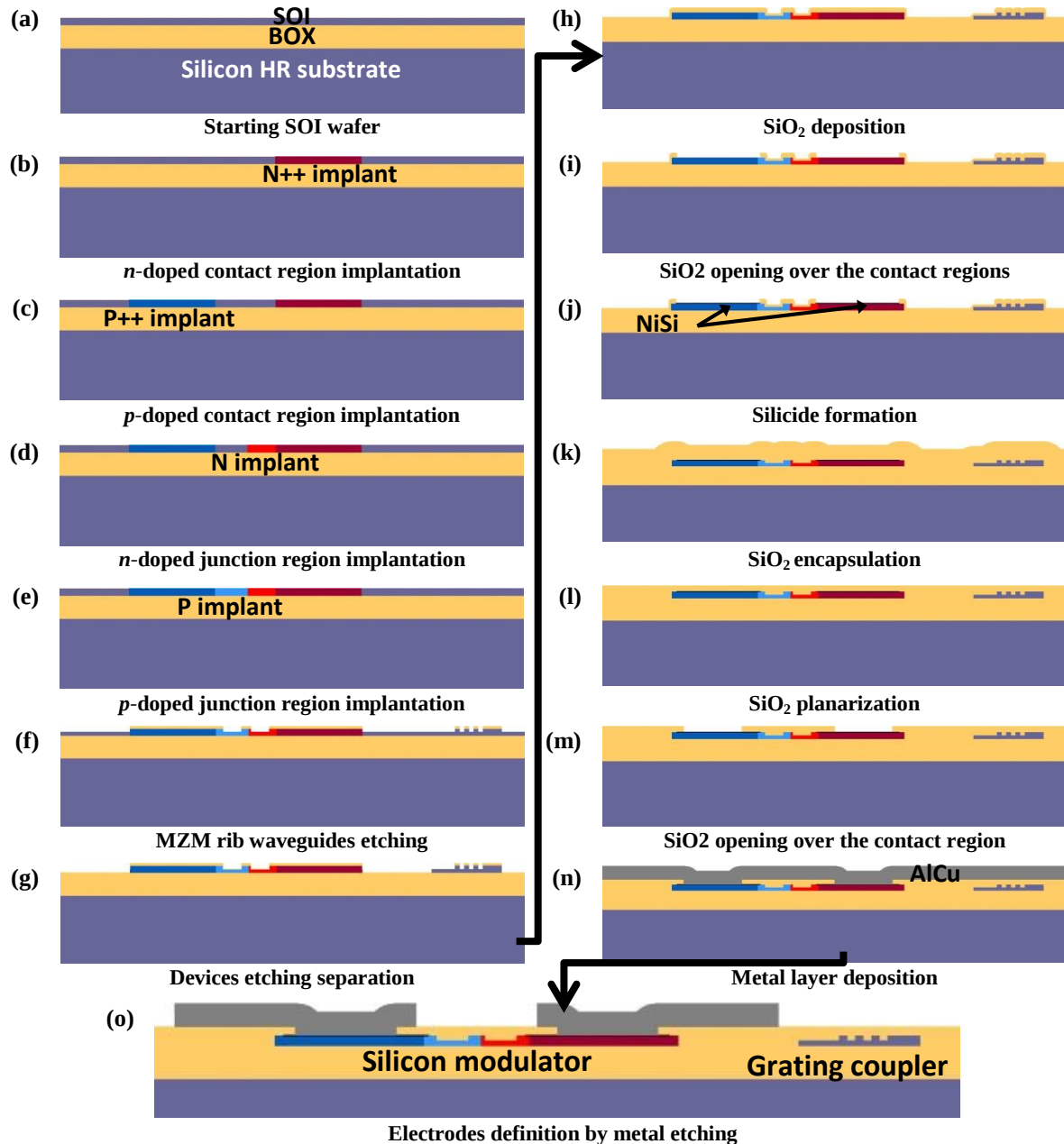


Figure II-22: Complete fabrication process flow of the silicon MZMs.

Ion implantation [see Figure II-22(a)-(e)]

The fabrication starts with a 5-nm-thick rapid thermal oxidation (RTO) at 1000°C during 55 seconds. The implantations steps are realized before etching the silicon waveguides, when the SOI surface is planar. Four regions are locally implanted to form the active regions of the MZMs. The implantation of each region (contact and junction) is realized using the steps described in **Figure II-23**.

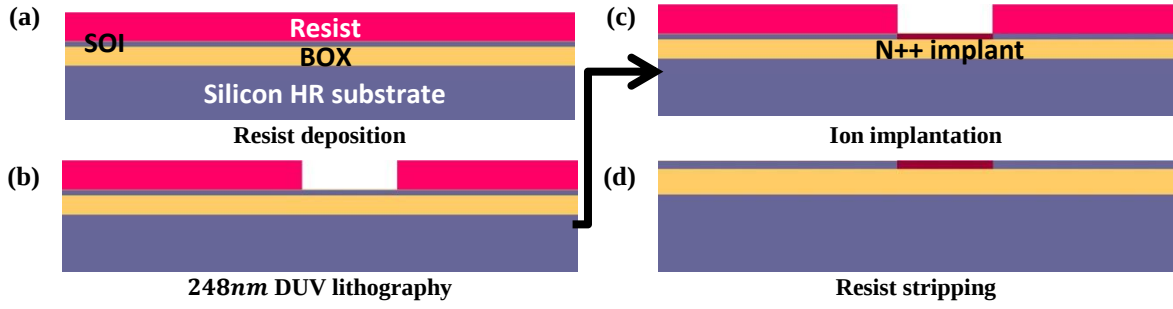


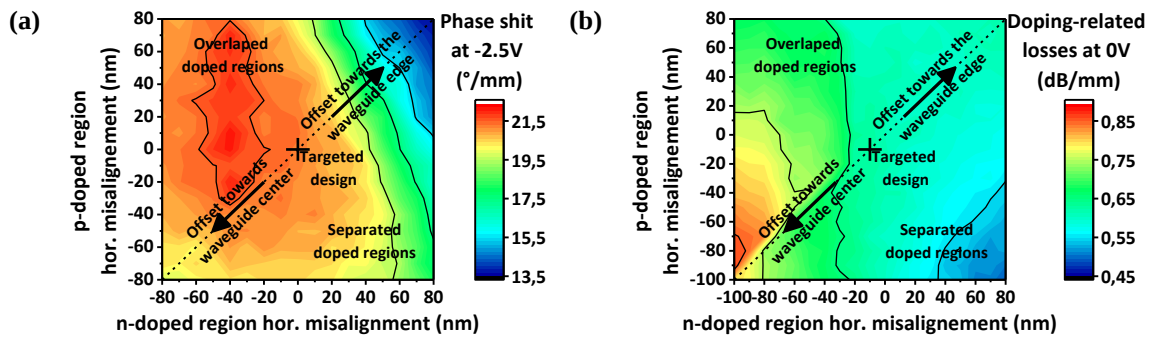
Figure II-23: Implantation steps for Si-doping

248nm deep ultraviolet (DUV) lithography is used for each implantation step. The exact doses for the p - n junctions, obtained by multiplying by 4 (for the p -type doping) and by 6 (for the n -type doping) the values proposed in **Table I-1** could not be used directly during the fabrication process, and have been slightly modified as shown in **Table II-4**. According to our simulations, these doses should give the same performances as displayed on **Figure II-13**. Double-step implantations are used for the contact regions, with an expected doping concentration of $2 \times 10^{19} \text{cm}^{-3}$ for both n -doped and p -doped contact regions. Boron and Phosphorus are respectively used for p -type and n -type doping.

Table II-4: Doping conditions used for the modulator.

Region	Doping type	Dopant	Tilt (°)	Twist (°)	Number of rotations	Dose (at/cm ²)	Energy (keV)
Contact region	<i>n</i> -doping	Phosphorus	7	27	4	4.8×10^{14}	180
						2.3×10^{14}	50
	<i>p</i> -doping	Boron				4.8×10^{14}	70
						2.3×10^{14}	20
<i>p</i> - <i>n</i> junction	<i>n</i> -doping	Phosphorus				1.5×10^{13}	200
						7×10^{12}	100
						3.4×10^{12}	35
	<i>p</i> -doping	Boron				8.7×10^{12}	80
			3.5×10^{12}	35			
			1.5×10^{12}	10			

During the lithographic steps of the implantation, the patterns are illuminated die by die, and might be misaligned from die to die. This misalignment can go up to $\pm 80 \text{nm}$ from the targeted position of the patterns. The effect of this misalignment on the p - n junction phase shift and DRL was evaluated using the model described in section II-2.1, and is displayed on **Figure II-24**. It can be seen that the phase shift at -2.5V can vary from $13.5^\circ/\text{mm}$ to $22^\circ/\text{mm}$, and the DRL from $0.5 \text{dB}/\text{mm}$ to $0.85 \text{dB}/\text{mm}$.

Figure II-24: Evolution of the (a) phase shift at -2.5V bias and (b) doping-related losses at 0V bias, versus lithographic misalignment of the doped p and n regions from their targeted positions.

Silicon waveguide patterning [see Figure II-22(f)-(g)]

The typical fabrication process used for the patterning of the silicon waveguides is depicted on **Figure II-25**. At first, a 100-nm-thick SiO₂ hard mask is deposited. The deposition is followed by 193nm DUV photolithography, hard mask reactive ion etching (RIE), resist stripping, and 150-nm-deep silicon RIE. These steps are necessary to pattern the MZMs waveguides, MMIs and the waveguide-to-fiber grating couplers. It can be noted that a bottom anti-reflective coating (BARC) layer (not shown on **Figure II-25**) is deposited before the resist, in order to reduce the reflections from the wafer during the photolithography, and is etched before the hard mask. The BARC layer will be present for any DUV photolithography used to etch the silicon layer. The silicon patterning is followed by a 5-nm-thick RTO at 1000°C, in order to reduce the waveguides roughness.

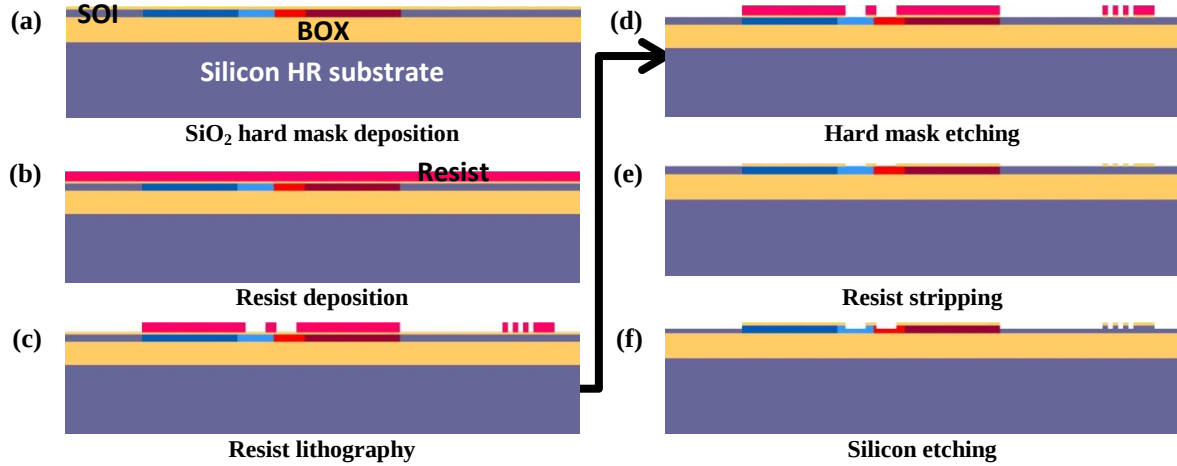


Figure II-25: Silicon waveguide patterning.

All devices are then separated using a similar process to the one depicted on **Figure II-25**, but without hard mask deposition. Here, 248nm DUV photolithography is used, directly followed by RIE of the left silicon layer up to the BOX, and resist stripping. SEM images of the different elements composing the MZMs are displayed on **Figure II-26**. After the silicon etching, the wafers were annealed in a N₂ environment during 15 seconds at 1050°C, in order to activate the implanted ions, even if the 1000°C RTO used to reduce the waveguides roughness might have been enough to do so.

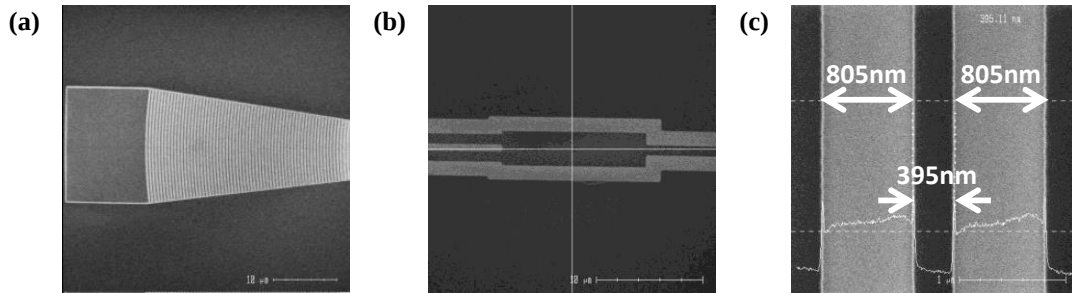


Figure II-26: SEM views of the (a) grating coupler, (b) MMI, and silicon waveguide forming the (c) MZMs arms.

Silicide formation [see Figure II-22(h)-(j)]

In order to reduce the resistance of the contact regions and limit their impact on the access resistance, these specific areas were silicided. At first, a 100-nm-thick SiO₂ layer is deposited to protect the non-silicided areas. The specific areas which need to be silicided are opened using 248nm DUV lithography and SiO₂ RIE, as shown on **Figure II-27(a)**. Then, the wafers are chemically cleaned using Hydrofluoric acid (HF), before an in-situ dry etching of the native oxide and the deposition of a bi-layer made of Ni (9nm) and TiN (10nm) (see **Figure II-27(b)**). It is essential to remove any native oxide between the silicon and the Ni, or the silicide process will fail. The deposition is followed by an annealing in a N₂ environment during 30 seconds at 300°C, in order

to form NiSi_2 on the contact regions. The TiN and non-reacted Ni layers are selectively removed by chemical etching. Finally the wafers are annealed again in a N_2 environment during 30 seconds at 450°C to form a NiSi layer ($\approx 18\text{-nm}$ -thick) (see **Figure II-27(c)**). SEM images of the silicided regions are shown on **Figure II-28**.

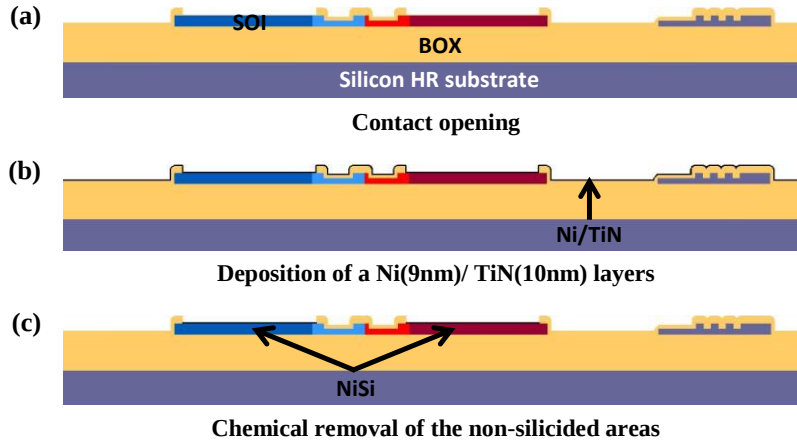


Figure II-27: Silicide formation.

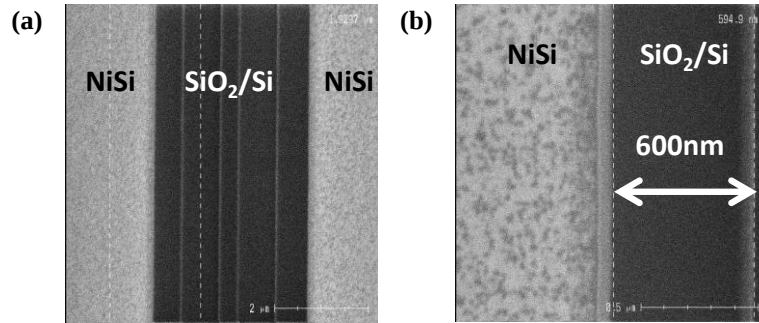


Figure II-28: SEM views of the MZMs silicided regions.

Planarization and electrodes patterning [see **Figure II-22(k)-(o)**]

Once the contact regions have been silicided, the wafers are encapsulated in 800nm of SiO_2 , and planarized by CMP. Approximately 300nm of SiO_2 is left above the waveguides. The silicided contact regions are re-opened using 248nm DUV lithography and SiO_2 RIE, as shown on **Figure II-29(a)**. A 650-nm -thick layer of AlCu is then deposited on the wafer (see **Figure II-29(b)**). Finally, the electrodes are patterned by using 248nm DUV lithography and RIE to etch the metal down to the SiO_2 , as shown on **Figure II-29(c) to (f)**.

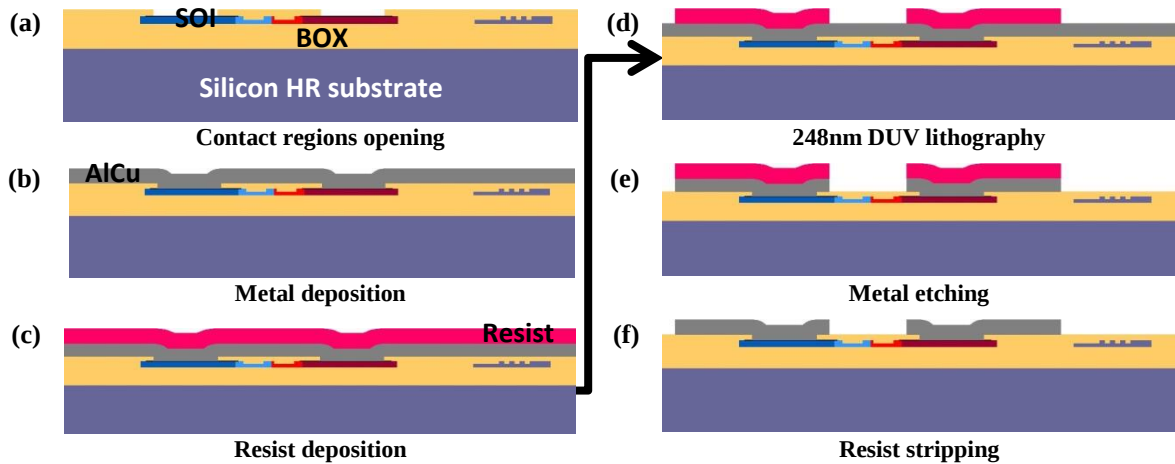


Figure II-29: Metal deposition and electrodes patterning

II-4. Mach-Zehnder modulator characterization

In this section, the characterization of the fabricated stand-alone modulators is detailed. First, the components forming the silicon MZMs (grating couplers, passive waveguides, MMIs, and bends) are optically tested, as well as the complete structures, to evaluate their optical losses. The MZMs are then electrically tested to estimate their efficiency, followed by small-signal measurements to evaluate their operating speed. Finally, large-signal measurements (eye diagrams) are realized to study their overall performances at a 25Gb/s bit rate. All those measurements have been conducted on the same $8''$ SOI wafer, using an automated probing station.

II-4.1. Passive components optical characterization

Before testing the complete MZMs structure, it is necessary to measure separately the different components forming the MZMs, in order to determine the different sources of optical loss and especially those corresponding to the p - n junctions. To do so, the measurement set-up depicted on **Figure II-30** is used. A benchtop laser sends an optical signal over a large wavelength span around $1.3\mu\text{m}$. Since the grating couplers are polarization-dependent, a polarizer is necessary to maximize the power coupled into the circuit. The power sent in the input coupler is set at 0dBm . Light goes through the silicon circuit, is coupled out, and is detected by a power-meter. SMFs are used both at the input and output of the circuit, and are fixed with a nearly-vertical angle of 11.5° .

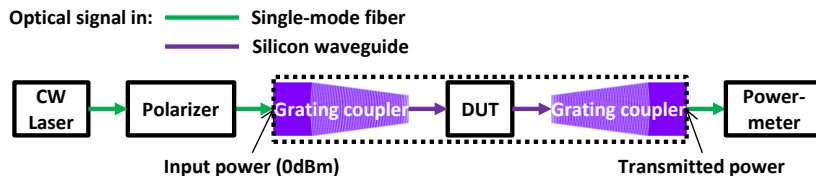


Figure II-30: Set-up for optical passive measurements.

The $8''$ SOI wafer is constituted of identical dies with dimensions of $22\text{mm} \times 22\text{mm}$, which gives more than 40 dies viable for testing. Automatic testing was used, where the fibers are first aligned above the grating couplers ($\approx 15\mu\text{m}$ gap with the wafer surface), before proceeding to the measurement itself. Due to time-limitation on the probing station use, the passive components have not been optically tested on all dies. Thus, 13 of them have been selected for these tests.

The first tested devices are the grating couplers. In this case, there is no device under test (DUT) as depicted on **Figure II-30**. The grating couplers are supposed to be identical, so we make the assumption that their transmission spectra will be the same. Since the input power is 0dBm , the optical loss of each coupler is half the power received by the power-meter. A raw measurement is presented on **Figure II-31(a)**, with interferences on the spectrum. We assume these interferences come either from reflection between the fibers and the grating couplers or from internal reflections between the gratings. Thus, a 2^{nd} order polynomial fit is performed on these raw data. The optical loss of the individual grating couplers on each die is shown on **Figure II-31(b)**. The variations between each die are related to fabrication variations, but also to the measurements. For instance, the gap between the fibers and the wafer surface was not monitored. Due to the bow of the $8''$ wafer, this gap was not the same between the dies at the center and those at the edge. Therefore, for the optical measurements, the mean value is taken as the reference, and the grating coupler loss (which corresponds to the maximum of transmission) is estimated at 4.2dB .

Using the same methodology, the performances of the MMIs were evaluated. In order to verify if their input power is effectively split in two, a single MMI is used as a DUT, with a grating coupler at each output. The average power difference between the two outputs is 0.1dB , which means that the splitting is $\approx 49/51\%$. By cascading two MMIs and measuring the power output difference between the two, the mean loss of the MMIs is estimated at 0.1dB .

The loss of the passive waveguides is evaluated by using two spirals of different lengths between the grating couplers. The difference of output power between the spirals, divided by the length difference gives an average waveguide loss of 0.8dB/mm , which is probably the result of a relatively large roughness on the edges of the rib waveguides. The loss corresponding to the 90° bends (with a radius of curvature of $30\mu\text{m}$) is also evaluated in the same way, by using two DUTs with a different number of bends. The loss per bend is evaluated at 0.06dB .

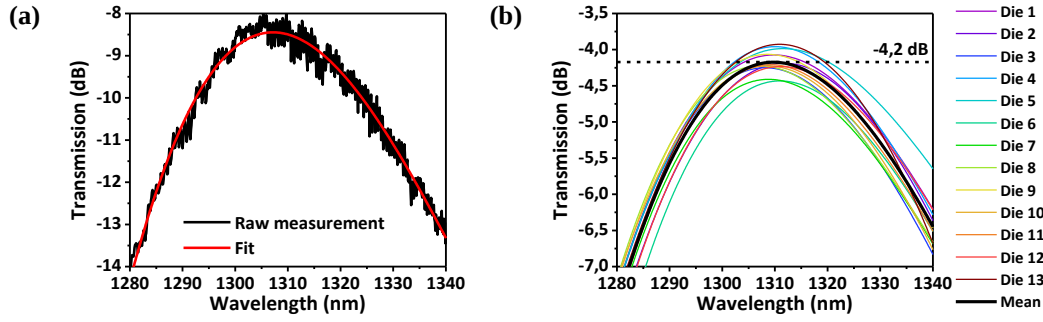


Figure II-31: (a) Example of raw measurement on a single die of the power received by the power-meter (input power: 0dBm). (b) Optical loss of the individual grating couplers on each die (fitted).

Finally the complete structures shown in **Figure II-21** are optically tested for the three different active regions lengths (2, 4, and 6mm). For these lengths, the average loss is respectively 15.4dB, 18.4dB and 21.4dB. Thus we can deduce that the total loss of the doped waveguide is 1.5dB/mm. Since the loss of the undoped waveguides is estimated at 0.8dB/mm, the DRL of the junction is estimated at 0.7dB/mm, which is close to the 0.63dB/mm obtained in the simulations shown on **Figure II-13(b)**. The results of the different measurements are summarized in **Table II-5**. It can be seen that the sum of optical losses of each component is coherent with the optical loss of the complete structure. It can be seen that the global performance of the passive components are quite far from the state-of-the-art of components with the same SOI thickness (especially the grating couplers and the passive waveguides) [11]. Thus, for a next demonstration, the grating couplers design should be improved, as well as the silicon waveguides edge roughness during the fabrication process.

Table II-5: Measured optical losses for each component, and in the complete MZM structures.

Component	Individual losses		MZM structure constitution		Total loss for each component (dB)
	dB/per component	dB/mm	Number of components	Length (mm)	
Grating coupler	4.2		2		8.4
90° bend	0.06		14		0.84
Passive waveguides		0.8		3.25	2.6
MMIs	0.1		2		0.2
Doped waveguides		1.5		2 / 4 / 6	3 / 6 / 9
Total structure loss (sum of components, for the 2 / 4 / 6mm-long active region)					15,04 / 18,04 / 21,04
Total structure loss (mean of measurements, for the 2 / 4 / 6mm-long active region)					15,4 / 18,4 / 21,4

II-4.2. Active region static electro-optical characterization

The next step is to measure the efficiency of the MZMs. The set-up used for these measurements is displayed on **Figure II-32**. This set-up is similar to the one of **Figure II-30 (b)**, but additionally, DC sources are used to electrically control the arms of the MZM. The central *p*-doped region is always connected to the ground. When one arm is under test, its *n*-doped region receives a positive potential, while the second *n*-doped region is connected to the ground. This way, there are no floating potentials on the electrodes when one of the arms is under test.

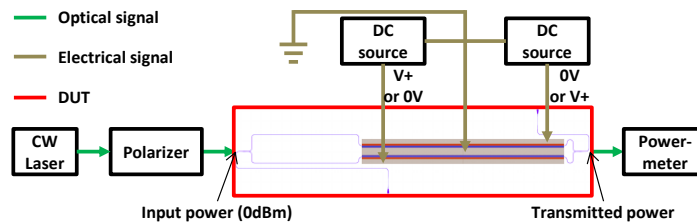


Figure II-32: Set-up for static electro-optical measurements.

The phase shift with the applied voltage on the junction is evaluated by observing the wavelength shift ($\Delta\lambda$) of the interferometric peaks when a bias is applied on the junction, and comparing it to the FSR of the spectrum, which is the wavelength spacing between two interferometric peaks. A wavelength shift of one FSR is equivalent to a 2π phase shift. Thus, the linear phase shift is expressed as:

$$\Delta\varphi(V_{bias})[^\circ/mm] = \frac{360 \cdot \Delta\lambda(V_{bias})}{L_m \cdot FSR} \quad (II.30)$$

The FSR is also used to determine the optical group index (n_g) of the silicon waveguides using [134]:

$$n_g = \frac{\lambda^2}{FSR \cdot \Delta L} \quad (II.31)$$

Thus, with an average $FSR \approx 4.2nm$, $\Delta L = 100\mu m$ (which is the fixed length difference between the MZM arms) and $\lambda = 1.3\mu m$, we obtain $n_g \approx 4$.

The average phase shift efficiency on each arm of the three MZMs was measured on several dies over the 8" SOI wafer. At a 2.5V reverse-bias, the phase shift over the 8" SOI wafer ranged from $13^\circ/mm$ to $19^\circ/mm$. We assume these variations are caused by the misalignments during the implantations of the $p-n$ junctions as explained in section II-3. Amongst the measured dies, the ones with the best efficiencies were selected to pursue the measurements. The spectra and phase shift efficiency of one of these dies are displayed on **Figure II-33**. It can be seen that the phase shift is relatively close to the one obtained by simulation.

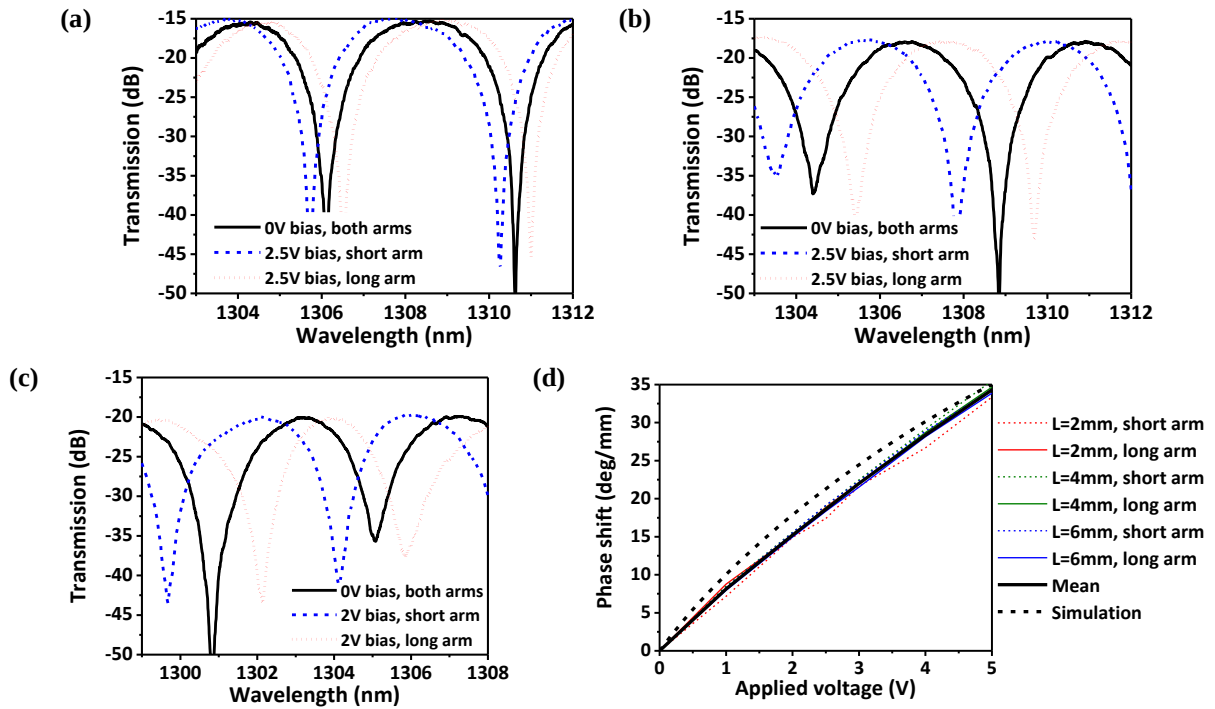


Figure II-33: Transmission spectra of the device under different bias from input fiber to output fiber, for (a) 2-*mm*-long, (b) 4-*mm*-long, and (c) 6-*mm*-long silicon MZM. (d) Phase shift versus applied voltage, simulations and measurements for all lengths.

II-4.3. Small-signal electrical and electro-optical characterization

Once the optical losses and efficiency of the modulator have been evaluated, the next step is to measure the speed of the MZMs. At first, the S-parameters are measured to validate the design of the electrodes, and followed by the modulation depth. The set-up used for the S-parameters measurements is displayed on **Figure II-34**. The MZMs are connected to an Agilent vector network analyzer (VNA), using two RF probes with a 100- μm -pitch. The impedance of each VNA port is at 50Ω . The S-parameters magnitude and phase are measured from 100MHz to 40GHz. A bias-Tee is also used to add a static signal, so the S-parameters can be evaluated at

several biases. To get rid of the influence from the coaxial cables, RF probes and bias tee, a 2-ports short-open-load-through (SOLT) calibration on an impedance standard substrate is realized before the measurements.

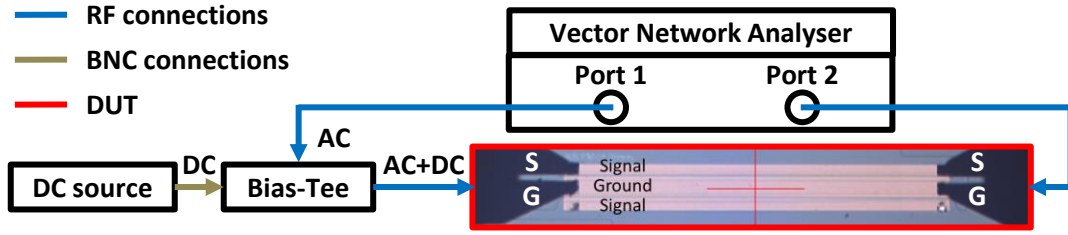


Figure II-34: Set-up for S-parameter measurements.

The results of these measurements for each MZM length are compared to the simulations shown in section II-2.2, and displayed on Figure II-35. For all lengths, the measured S_{11} parameters are in good agreement with the simulations. However, the measured S_{21} parameters -6dB electrical bandwidths are lower than predicted. The results are summarized in Table II-6.

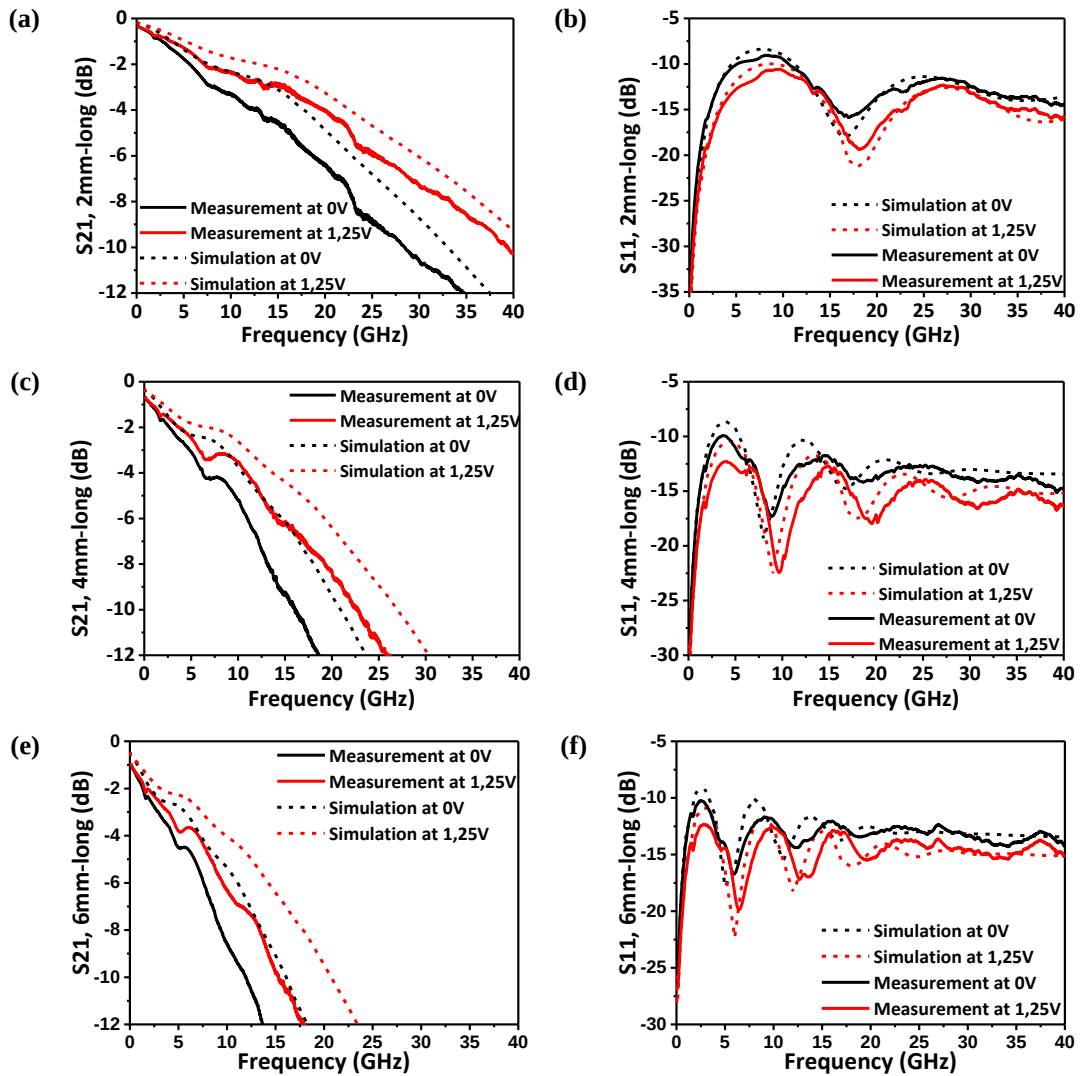
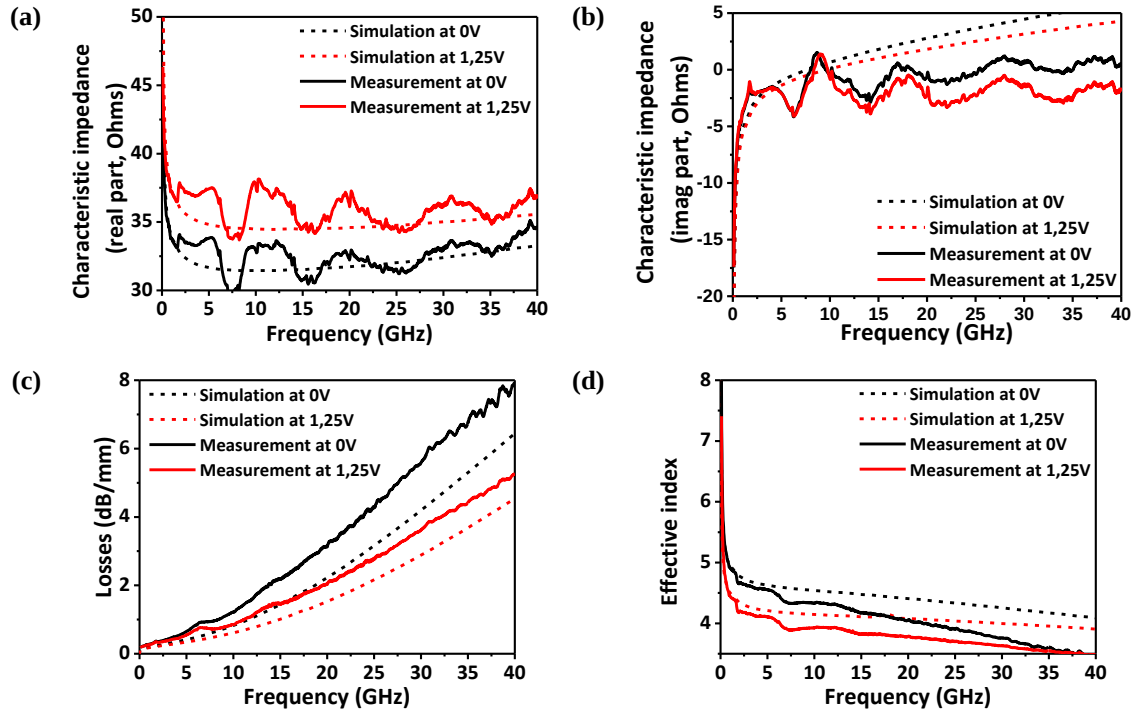


Figure II-35: Simulations and measurements of the S-parameters magnitudes for the different lengths of the active region: (a) and (b) 2-mm-long, (c) and (d) 4-mm-long, (e) and (f) 6-mm-long.

Table II-6: Simulated and measured S_{21} –6dB bandwidths.

Active region length (mm)	Simulated S_{21} –6dB bandwidth (GHz)		Measured S_{21} –6dB bandwidth (GHz)	
	0V bias	1.25V reverse-bias	0V bias	1.25V reverse-bias
2	22.9	29.8	18.7	25.4
4	14.7	19.2	11	14.1
6	11.1	14.2	7.6	9.5

To understand the differences between simulations and measurements, the characteristic impedance, RF losses and RF effective index are extracted from the S-parameters of the 4-mm-long MZM – using the conversion formulas in [138] –, and compared to the simulations, as depicted on **Figure II-36**. While the characteristic impedance is relatively close to the simulations (the variations being due to measurement artefacts), the RF losses are higher and the RF effective index is lower than predicted by the simulations, which explains the bandwidth reduction. Two possibilities can explain these differences. The first one is the substrate resistivity, which is unknown after the fabrication process of the MZM, and may have been degraded by parasitic surface conduction [141]. The second one is due to the real geometry of the electrodes, which is also unknown, and may have changed the resistance and the inductance of the line.

**Figure II-36:** Simulations and measurements of the (a) real and (b) imaginary part of the characteristic impedance, and of the (c) RF losses and (d) RF effective index.

Finally, the modulation depth of the MZMs is measured by using the set-up displayed on **Figure II-37**. One port of the VNA is connected at the beginning of the line, as in the previous set-up, but the end of the line is terminated with a DC-block and a load with an impedance of 50Ω . A constant optical signal is sent in the MZMs, modulated by the MZMs, converted in an electrical signal by an Agilent N4373C Lightwave Component Analyzer (LCA), and sent back to the second port of the VNA. By normalizing the small-signal frequency response to the lowest output frequency of the VNA, we obtain the modulation depth.

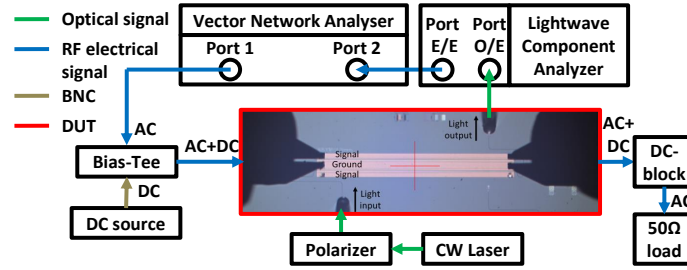


Figure II-37: Set-up for small-signal electro-optical measurements.

The measured modulation depths for the different lengths are shown on **Figure II-38**. The large level of noise beyond 20GHz comes from the low level of signal on the LCA, since no amplification system is used during the measurements. The E/O bandwidths are summarized in **Table II-6**. It can be seen that for the 4 and 6-mm-long MZM, the measurement get closer to the simulation, even though the RF losses are higher. This may be due to an improved matching between the RF effective index and optical group index.

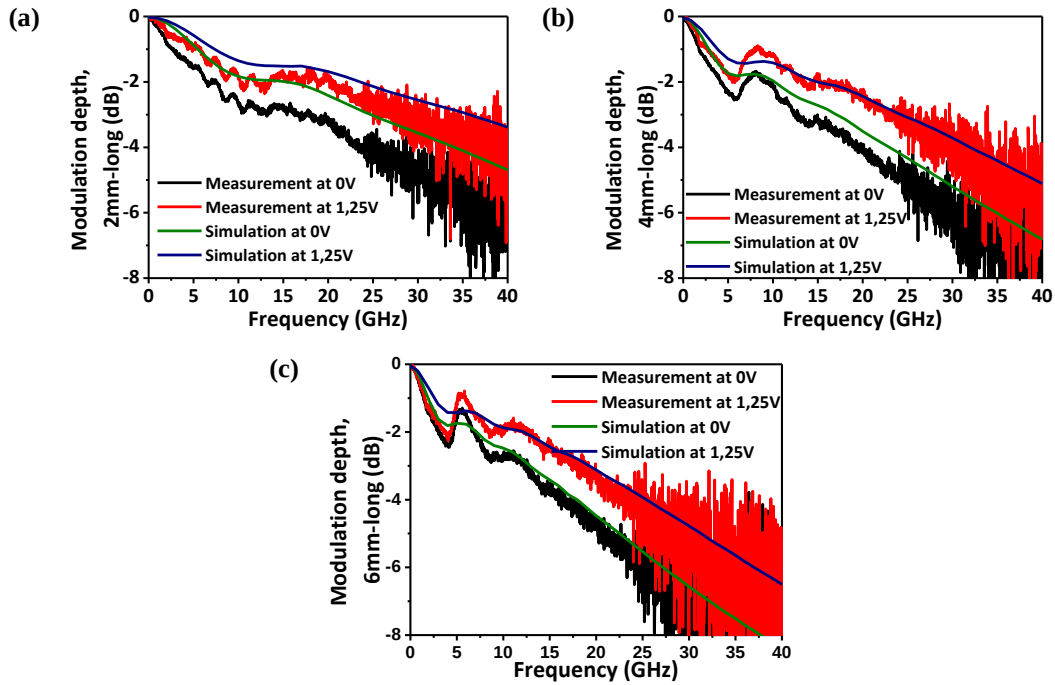


Figure II-38: Simulations and measurements of the modulation depth for the different lengths of the active region: (a) 2-mm-long, (b) 4-mm-long, and (c) 6-mm-long.

Table II-7: Simulated and measured electro-optical bandwidths.

Active region length (mm)	Simulated electro-optical bandwidth (GHz)		Measured electro-optical bandwidth (GHz)	
	0V bias	1.25V reverse-bias	0V bias	1.25V reverse-bias
2	24.7	35.7	17.3	28.2
4	17.1	24.3	12.7	23.1
6	12.9	19.3	12.3	18

II-4.4. Large-signal electro-optical characterization^d

Before realising the large-signal E/O measurements, it is necessary to fix the configuration in which the MZMs must be used. To help us to do so, a model of the silicon MZMs was realized in the circuit simulator of PICWAVE, which takes into account the passive and the static electro-optical measurements of the previous

^d The large-signal electro-optical measurements being very time-consuming, we did not have time to measure all three MZMs with their different lengths. Therefore, due to their high level of losses and lower speed, the 6-mm-long MZMs were not studied any further.

sections. Using this model, we were able to reproduce the optical spectra of the silicon MZM, as shown on **Figure II-39**. This model will also be used later in this section, to simulate eye diagrams.

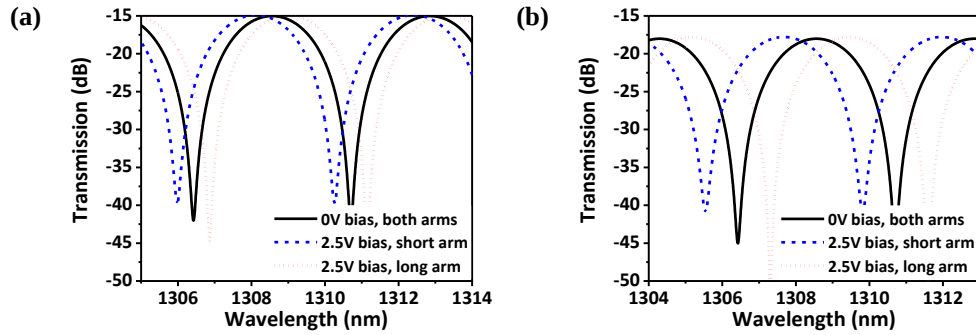


Figure II-39: Simulation of the transmission spectra of the device under different bias from input fiber to output fiber, for (a) 2-mm-long, and (b) 4-mm-long silicon MZM.

Since the MZMs are based on a SGS configuration, both arms can be controlled separately to enhance the efficiency for a given voltage sweep. This mode of operation is referred as dual-drive push-pull. By alternatively applying 0V or $-2.5V$ on the arms of the MZM (as shown on **Figure II-40**), the ER will be higher than by using a single arm. In the following measurements, both arms will receive the same voltage sweep ($2.5V_{pp}$), and the same fixed bias ($-1.25V$), so that the voltage sweep on the junction will vary from 0 to $-2.5V$.

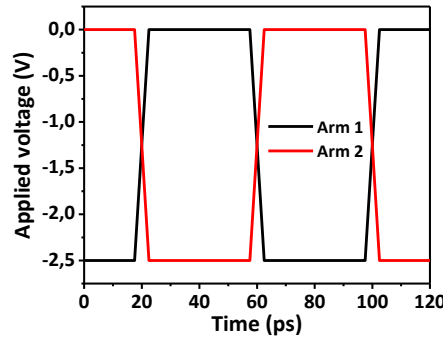


Figure II-40: Typical form of the electrical signals applied on each MZMs arm, in a dual-drive push-pull operation at 25Gb/s.

Even if the MZMs operates by applying an electrical signal on the arms, the static phase difference between the two arms – equivalent here to the wavelength of operation – must be carefully chosen to maximize the performances of the MZM. As it can be seen on **Figure II-41**, for MZMs operating in dual-drive push-pull operation, the ER will be null if the static phase difference between the arms is equal to 0 (maximum of transmission) or π (minimum of transmission). By increasing the static phase difference between 0 and π , the extinction ratio will increase, but at the cost of higher optical losses.

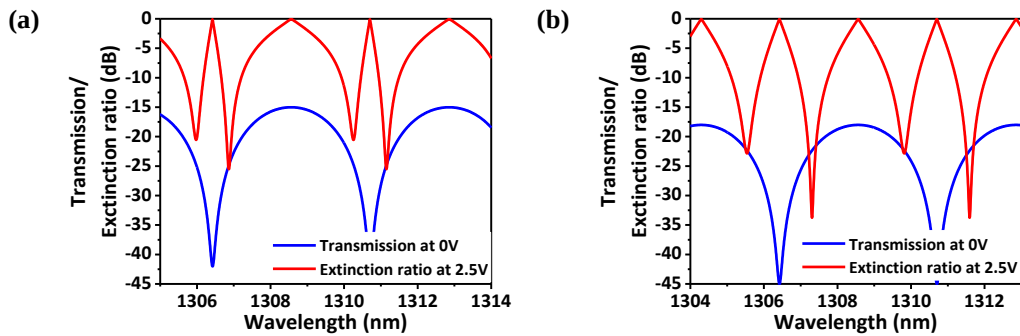


Figure II-41: Simulated transmission at 0V and ER when the MZM is operated in dual-drive push-pull, with $-2.5V$ applied on each arm for (a) a 2-mm-long active region and (b) a 4-mm-long active region.

When the static phase difference between the arms is equal to $\pi/2$ (or at -3dB from the maximum of transmission), the MZM is said to be at quadrature. As it is shown on **Figure II-42(a)-(b)**, the OMA is maximized when the MZMs are set at quadrature. Therefore, for the next measurements, the wavelength will be chosen so that the MZMs are set at quadrature.

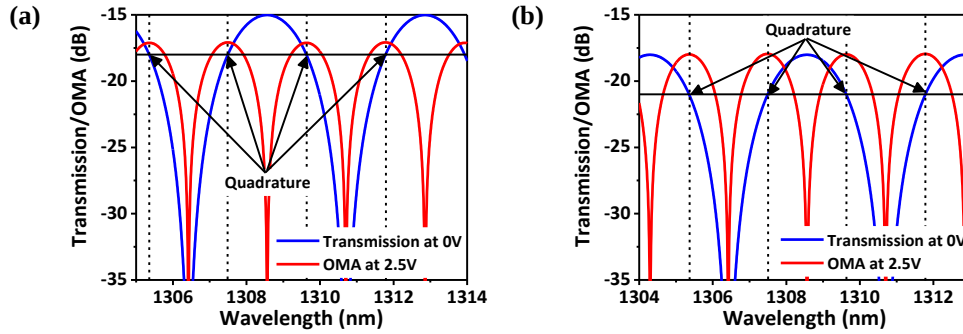


Figure II-42: Simulated transmission at 0V and OMA when the MZM is operated in dual-drive push-pull, with -2.5V applied on each arm for (a) a 2-mm-long active region and (b) a 4-mm-long active region.

The set-up for the large-signal E/O measurements is depicted in **Figure II-43**. The electrical driving signals are provided by a M8020A pattern generator associated with a M8061A Multiplexer 2:1 up to 32Gb/s , both from Keysight. Since the MZMs are driven in a dual-drive push-pull configuration, both the $DATA$ and $\overline{DATA}$ outputs are used, to send signals in opposite phase on each MZMs arm. Even so, the electrical signals pass through RF phase shifters, to ensure they are perfectly in opposite phase on the arms. The signals are then amplified using two paired SHF 806E amplifiers and sent to the MZMs via an SGS RF probe. Since the voltage sweeps sent by the generator are centered at 0V, DC sources and bias-Tees are used to offset the signals. The other end of the modulator is terminated with a DC-block and a 50Ω load.

Input light is provided by a tunable laser source with an output power of $+13\text{dBm}$. For all measurements, the laser wavelength is set at quadrature. Light at the device output passes through a booster optical amplifier (BOA) – which is polarization-dependent – and an optical tunable filter. This filter is used to reduce the noise from the BOA spontaneous emission. Finally, the signal is monitored using the high-speed optical entry of a Keysight 86100D DCA-X Oscilloscope.

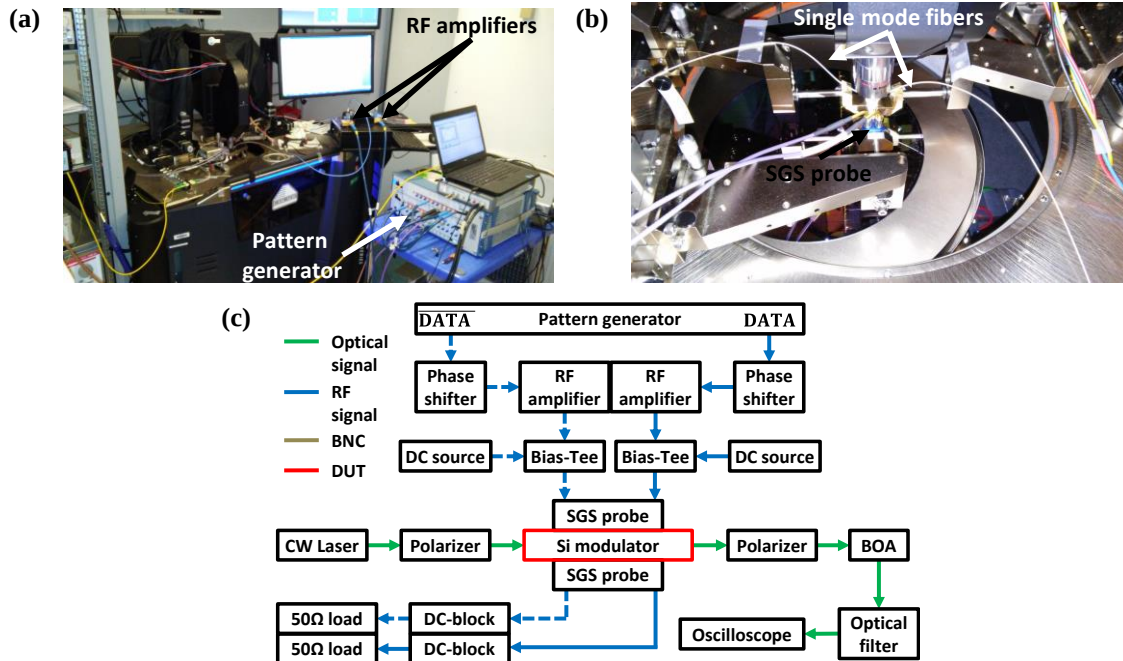


Figure II-43: (a) and (b) MZM under high-speed testing. (c) Set-up for large-signal electro-optical measurements.

While the best option would be not to use the BOA – since it brings a lot of noise to the measurement if its gain is too high – its use was made necessary because of the low sensitivity of the high-speed detector. A compromise is to drive the BOA with a low current, thus the combination of the BOA and the optical filter give a power amplification estimated between $+5$ and $+8\text{dB}$, which is enough to conduct our measurements.

In order to remove the influence of the meter-long RF cables used in the set-up (displayed on **Figure II-43**), electrical de-emphasis provided by the multiplexer is used. The de-emphasis is only used to de-embed the signal degradations caused by cables, and not to compensate the MZMs frequency response. An example of electrical signal before and after de-emphasis is shown on **Figure II-44**. It can be seen that the transition times between “ON” and “OFF” states are largely improved.

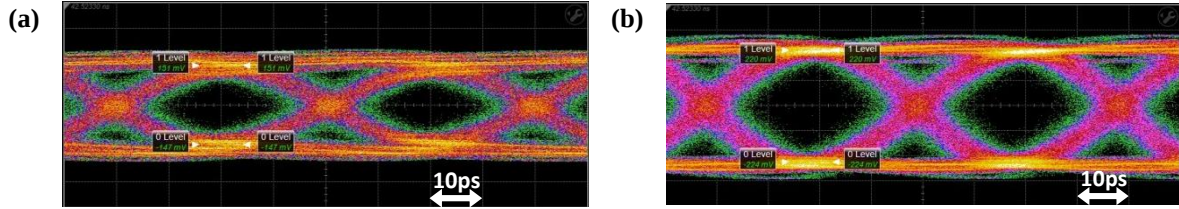


Figure II-44: Eye diagrams at 25Gb/s of the electrical signal send on the RF probes: (a) without de-emphasis, (b) with de-emphasis.

The measured eye diagrams are shown on **Figure II-45(a)** and **(c)**. Both eyes are open at 25Gb/s , in a dual drive configuration, with $2.5V_{pp}$ send on each arm, biased at a 1.25V reverse-bias. A passive numerical low-pass filter (4th order Butterworth filter) with a 25GHz cut-off frequency was also used to reduce the noise during the measurements. The measured extinction ratios are respectively 3.9dB and 7.1dB for the 2-mm -long and 4-mm -long MZM.

The RF parameters of the line extracted in the previous section were also added to the PICWAVE model to realize eye diagram simulations. The resulting eye diagrams are shown on **Figure II-45(b)** and **(d)**, and show a good agreement with the measured eye diagrams. By using the simulations, we were also able to evaluate more precisely the insertion loss of the MZMs at 25Gb/s . The performances of the measured MZMs are summarized in **Table II-8**.

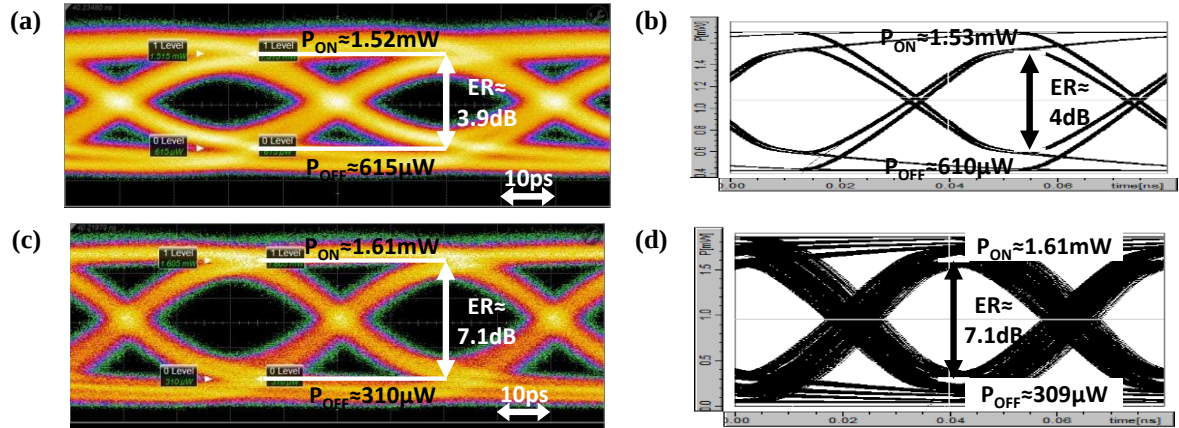


Figure II-45: Measured and simulated eye diagrams at 25Gb/s of the (a,b) 2-mm -long and (c,d) 4-mm -long silicon MZM, in dual drive configuration with $2.5V_{pp}$ applied on each arm, each biased at -1.25V .

Table II-8: Summarized performances of the MZMs from input SMF to output SMF.

Active region length (mm)	DRL (dB/mm)	Phase shift at -2.5V ($^\circ/\text{mm}$)	$V_\pi L_\pi$ at -2.5V (V.cm)	Maximum of transmission ^e (dB)	E/O bandwidth at -1.25V (GHz)	ER at 25Gb/s with $2.5V_{pp}$ (dB)	IL at 25Gb/s with $2.5V_{pp}$ (dB)
2	0.7	18.6	2.4	-15	28.2	3.9	16.4
4				-18	23.1	7.1	18.6

^e From input grating coupler to output grating coupler.

In order to evaluate the performances of our modulators, we tried to compare them with the current state-of-the-art silicon MZMs based on lateral p - n junctions. However, the performances of these modulators are displayed at their highest possible bit rates, which are generally above 40Gb/s , while we only performed measurements at 25Gb/s . Since we could not evaluate the performances of our modulators at higher bit rates (due to the limitations of our eye-diagrams testing set-up), a direct comparison is not possible. Nevertheless, by studying the different characteristics (efficiency, losses and bandwidth) of these modulators, an indirect comparison can be realized. These characteristics are summarized in **Table II-9**, with silicon MZMs driven on a single arm, or using both arms in single- or dual-drive push-pull configuration. Unlike the dual-drive configuration used for our modulators, in the single-drive case, both p - n junctions are directly connected and driven by the same SG electrodes. A static bias is applied between the diodes so that they are both reversely-biased. The main interest of this configuration is to reduce the total capacitance (since the p - n junctions are in series), and to simplify the driving circuit design. However, it is done at the cost of a higher driving voltage [126].

The DRL are rarely provided in the literature. Nevertheless, it can be seen that the propagation losses in our doped waveguides (1.5dB/mm) is comparable to those presented in **Table II-9**, even with the relatively high propagation losses in our waveguides. The same can be said for the $V_\pi L_\pi$ product of our modulator. As for the optical losses of our modulator (from input MMI to output MMI), based on values displayed on **Table II-5**, they can be respectively estimated at 4.9dB for the 2-mm -long MZM and 7.9dB for the 4-mm -long MZM, which are also comparable but slightly higher than the state-of-the-art. Therefore, except for the optical losses which should still be improved, the static performances of our modulators are within the norm. The E/O bandwidths of our modulators are honourable, but are still lower than the fastest devices with similar lengths ([126]–[128]), which had access to standard multi-level metallization for the designs of their electrodes.

Table II-9: State-of-the art of the silicon MZMs based on lateral p - n junctions, with different driving schemes.

Driving scheme	Active region length (mm)	Doped waveguide propagation losses (dB/mm)	$V_\pi L_\pi$ ^f (V.cm)	Stand-alone modulator optical losses ^g (dB)	E/O bandwidth (GHz)	Bias voltage (V)	Data rate (Gb/s)	Voltage sweep (V_{pp})	ER (dB)	Ref.
Single arm	0.75	1.6	$1.6 - 2$ ($1 - 8V$)	1.9	27.8	-5	60	6.5	3.6^h	[123]
	1	2.8	$1.5 - 2$ ($1 - 8V$)	3.9	30	-2	44	5	2.4	[124]
	2			6.2	19	-2	35	3	5.84	
Single-drive push-pull	6	1.2	$1.8 - 2.1$ ($3 - 5V$)	9	20	-2	30	3.5	8.5	[125]
	4	Unknown	3.2 ($4V$)	3.8	41	-4	60	4.8	3.8^f	[126]
Dual-drive push-pull	3	1.2	2.3 ($8V$)	8	27	-1	40	1.6	4^f	[127]
	3	1.1	2.5 ($1V$)	5.5	30	0	50	1.5	3.4^f	[128]
	2	1.5	2.4	4.9	28	-1.25	25	2.5	3.9	This work
	4		($2.5V$)	7.9	23.1	-1.25			7.1	

Even if we were not able to evaluate the performances of our modulators above 25Gb/s , we could use the PICWAVE model to simulate the eye diagrams. The results are displayed on **Figure II-46** for a 40Gb/s data rate, and on **Figure II-47** for a 50Gb/s data rate. As it can be seen, both modulator lengths should be able to operate at 40Gb/s , but due to its lower E/O bandwidth, the 4-mm -long MZM might not be able to reach 50Gb/s . As expected, the ER would be lower than at 25Gb/s (while the IL would be higher), but would still compare honourably with the state-of-the-art. Therefore, the modulators performance should be more than enough for the integrated transmitters.

^f The reverse-bias voltage for the $V_\pi L_\pi$ values is given in parenthesis.

^g Generally, only the optical losses of the component from splitter to combiner are provided, without the grating couplers or the silicon waveguides outside of the MZM. The indicated losses correspond to the static maximum of transmission, and not the IL at a given bit rate.

^h Measurement at quadrature.

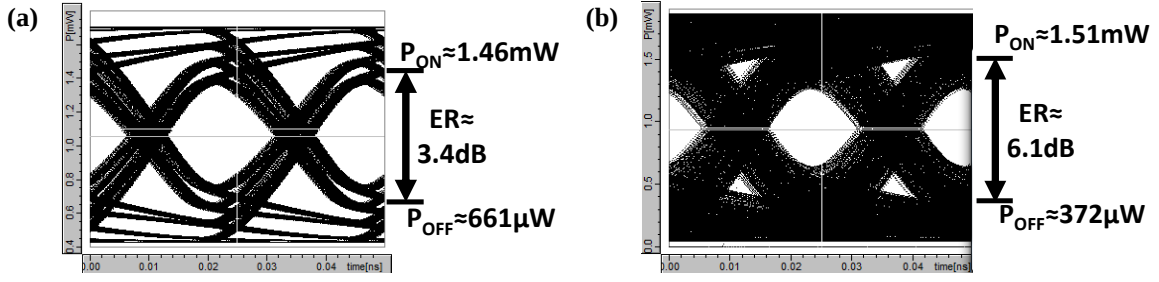


Figure II-46: Simulated eye diagrams at 40Gb/s of the (a) 2-mm-long and (d) 4-mm-long silicon MZM, both in dual drive configuration with $2.5V_{pp}$ applied on each arm, each biased at $-1.25V$.

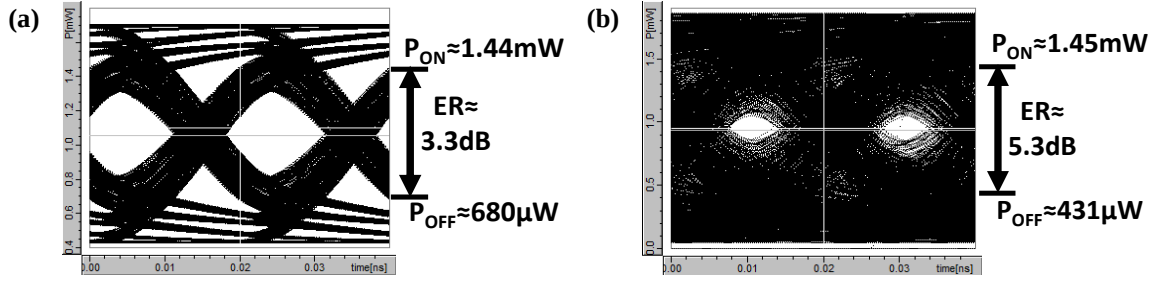


Figure II-47: Simulated eye diagrams at 50Gb/s of the (a) 2-mm-long and (d) 4-mm-long silicon MZM, both in dual drive configuration with $2.5V_{pp}$ applied on each arm, each biased at $-1.25V$.

Using the performances shown in **Table II-8** and **Eq. (II.10)**, it is now possible to evaluate *OMA* for the two MZMs, as a function of input power. The result is displayed on **Figure II-48(a)**. It can be seen that in the case of the fabricated devices, the MZM with a 2-mm-long active region is slightly better. However, due to the large optical losses of the passive components displayed on **Table II-5** (especially the grating couplers and the passive waveguides), more than 17.3dBm (or 53.7mW) of optical power would be required from an external laser to reach the targeted minimum *OMA* (the additional optical losses coming from a 10-km-long propagation in a SMF are not taken into account). Thus, in the case where the laser is integrated, the losses from the input grating coupler can be removed, and the laser should provide at least 13.1dBm (or 20.4mW) of optical power in the silicon waveguide (if the optical losses of the other components are not modified by the laser integration).

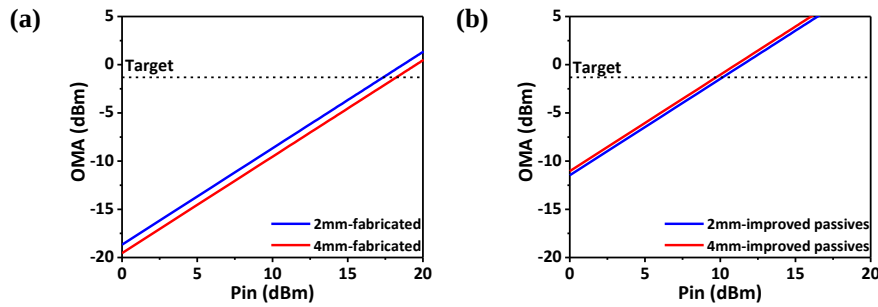


Figure II-48: Calculated *OMA* at 25Gb/s for (a) the fabricated MZMs, and (b) MZMs with the same *p-n* junction and TWE electrodes, but improved passive components. In both cases, the additional optical losses coming from a 10-km-long propagation in a SMF are not taken into account.

In order to evaluate the MZMs active regions with an improved grating coupler design and passive silicon waveguides with less surface roughness, we picked the performances of state-of-the-art passive components with the same SOI thickness [11], and used them in the PICWAVE model (2.2dB for the grating couplers, and 0.18dB/mm for the silicon waveguides). The expected performances and *OMAs* of these improved MZMs modulators are shown on **Table II-10** and **Figure II-48(b)**. It can be seen that in this case, the MZM with a 4-mm-long active region becomes better than the 2-mm-long one, and the required optical power from an external laser has been reduced to 9.7dBm (or 9.3mW). Thus, for an integrated laser, only 7.5dBm (or 5.6mW) of optical power should be coupled in the silicon waveguide to reach the targeted *OMA*.

Table II-10: Expected performances of the MZMs from input SMF to output SMF, with the same junction and TWE, but improved performances for the grating couplers and passive silicon waveguides.

Active region length (mm)	DRL (dB/mm)	Phase shift at -2.5V (°/mm)	E/O bandwidth at -1.25V (GHz)	Maximum of transmission (dB)	ER at 25Gb/s with 2.5V _{pp} (dB)	IL at 25Gb/s with 2.5V _{pp} (dB)
2	0.7	18.6	28.2	-7.8	3.9	9.2
4			23.1	-9.5	7.1	10.1

II-5. Conclusions and possible improvements

In this chapter, we presented the design, fabrication and characterization of silicon modulators which will be co-integrated with the hybrid III-V on silicon lasers. Amongst the several types of existing physical phenomena and device structures for optical modulation, we chose to rely on a Mach-Zehnder modulator based on carrier depletion in a reversely biased p - n junction, for its inherent tolerance towards fabrication variations and high-speed capabilities. While an EAM based on a hybrid III-V on silicon waveguide might also have been an interesting alternative, it would have been necessary to bond a different III-V stack from the one used for the laser, and also required more complex fabrication steps on the III-V waveguide.

The design of the p - n junction embedded in each MZM arms, and of the travelling-wave electrodes, has been detailed. The complete fabrication process, as well as the electro-optical characterization of the devices (through static, small-signal, and large-signal measurements) was also presented. We found a good agreement between the expected and effective performances of the devices, which comply with the requirements set in the **section II-1.5**: they operate with 2.5V_{pp} at 25Gb/s in the 1.3μm wavelength region, can be directly integrated with the hybrid III-V on silicon lasers (since they use only one metal level), and have been defined in a 300-nm-thick SOI layer (standard thickness used by the CEA-LETI and STMicroelectronics).

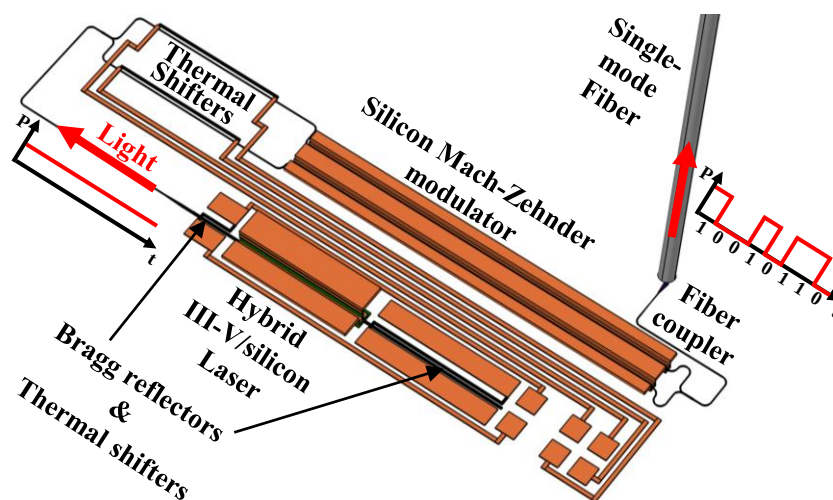
We also compared the performances of the designed modulators to the state-of-the-art. While a direct comparison was not possible due to the difference in bit rates, we still managed to realize an indirect comparison with the characteristics of each device. The static performances of our modulators are equivalent to those found in the state-of-the-art, even though they suffer from relatively large optical losses. We were able to determine that the source of the losses was mainly due to the passive components and not the active regions of the modulators. While being lower than the fastest devices, the E/O bandwidths of our modulators are still honourable, since we were limited to a single metal level to define the TWE. We also simulated their performances at higher bit rates, and showed that both 2-mm-long and 4-mm-long MZMs might operate at 40Gb/s, and even at 50Gb/s for the 2-mm-long one.

Based on the characterization of our modulators, we could evaluate the *OMA* for different input power from the laser source. Given the optical losses of our devices, a large level of power would be required to reach the targeted value of -1.3dBm. Even if the laser is heterogeneously integrated, more than 13dBm (or 20mW) should be coupled in the silicon waveguide given the current performances of the modulators (by supposing they would remain unchanged by the laser integration), without taking into account the additional optical losses coming from a 10-km-long transmission. Nevertheless, by improving the design of the grating coupler and the roughness on the edges of the silicon waveguides, in order to improve their losses, the necessary optical power coupled into the silicon waveguide would be reduced to 7.5dBm (or 5.6mW).

While the performances of the MZMs are sufficient for the transmitter, the study also enables us to find the possible ways to further improve the MZMs design. Thus, it would be useful to improve the trade-off between characteristic impedance, RF losses and RF effective index, and therefore enhance the bandwidth. For instance, by using more advanced design for the electrodes, such as segmented (or slow-wave) electrodes, the characteristic impedance can become closer to 50Ω [126], [127], [142], [143]. Additionally, the p - n junction doping concentrations were chosen as a trade-off between losses and efficiency, but one should also take into account the junction capacitance for the travelling-wave electrodes. Even if it means to reduce the phase shift, reducing the doping concentration would not only reduce the optical losses, but also the junction capacitance, which would simplify the trade-off between Z_c and $\tilde{n}_{RF}$, as explained in section II-2.2. It would also reduce the conductance between the lines, thus the RF losses (see **Eqs. (II.23)** and **(II.26)**). The larger access resistance resulting for the lower doping concentrations in the junction could be compensated by using intermediary doping levels between the junction and the contact regions [126]–[128].

Chapter III: Integrated III-V on silicon high-speed transmitter

This third chapter is dedicated to the demonstration of an integrated III-V on silicon high-speed transmitter for silicon photonics application in the data-centers. This transmitter makes use of the silicon Mach-Zehnder modulator demonstrated in the previous chapter, which is co-integrated with a hybrid III-V on silicon laser. At first, the design of the laser is presented, followed by the structure of the complete transmitter. Then, the fabrication process of the transmitter is detailed with the steps before and after bonding of the III-V wafer. Finally, the several components composing the transmitter are characterized before the complete structure testing, in order to fully understand the performances of the transmitters.



III-1. Hybrid III-V on silicon DBR laser	64
III-1.1. Laser overall structure view	64
III-1.2. Adiabatic taper design	66
III-1.3. Reflectors design	70
III-2. Transmitter structure	74
III-2.1. Silicon thickness transition.....	74
III-2.2. Silicon modulator	75
III-2.3. Transmitter layout	76
III-3. Hybrid III-V on silicon transmitter fabrication	77
III-3.1. SOI part fabrication	77
III-3.2. III-V integration.....	79
III-4. Integrated transmitter characterization	84
III-4.1. Stand-alone components.....	84
III-4.2. Complete transmitter characterization.....	89
III-5. Conclusions and improvements	93

III-1. Hybrid III-V on silicon DBR laser

Amongst the different kinds of hybrid III-V on laser sources presented in **section I-3.3**, the DBR structure was chosen for the integrated III-V on silicon transmitter. This structure is chosen for being able to operate in the single-wavelength regime, and more importantly for its wavelength tuning capabilities. The conception of this DBR laser is presented in this section, starting from a description of the overall structure, and followed by the designs of the elements composing the laser: the tapers used for adiabatic coupling and the reflectors in the silicon waveguide.

III-1.1. Laser overall structure viewⁱ

Schematic views of the hybrid III-V/silicon laser can be seen in **Figure III-1**. The bonded III-V stack is positioned above the silicon waveguide to provide the optical gain, with a 100-nm-thick SiO₂ gap between the two high refractive index media.

III-V waveguide description

The laser is based on the popular InGaAsP/InP system, which is largely used for long-distance optics when processed on their base InP wafers. The gain is provided by intrinsic InGaAsP multiple quantum-well (MQW) layers and barriers, rather than a simple double heterostructure or a simple quantum-well, in order to reduce the threshold current while keeping a good confinement in the gain region [3]. The gain spectral distribution of the MQWs depends on their composition and their amount of strain. In our case, the MQWs maximum gain is centered on 1.3 μ m. They are surrounded by two separate-confinement heterostructure (SCH) layers, with a lower refractive index than the MQWs to further improve the optical mode confinement in the gain region. The SCH layers are enclosed by two *p*- and *n*-doped InP layers, to form a *p-i-n* junction for electrical injection. The bandgap of the InP layers being larger than the bandgap of the InGaAsP layers, there is no interband absorption in these layers, and they will offer a good electrical confinement for the carriers injected in the active gain region, thus improving their radiative recombination probability. The *n*-doped layer is placed between the III-V and SOI waveguides rather than the *p*-doped layer. The reason for this choice is explained in **section III-3.2**.

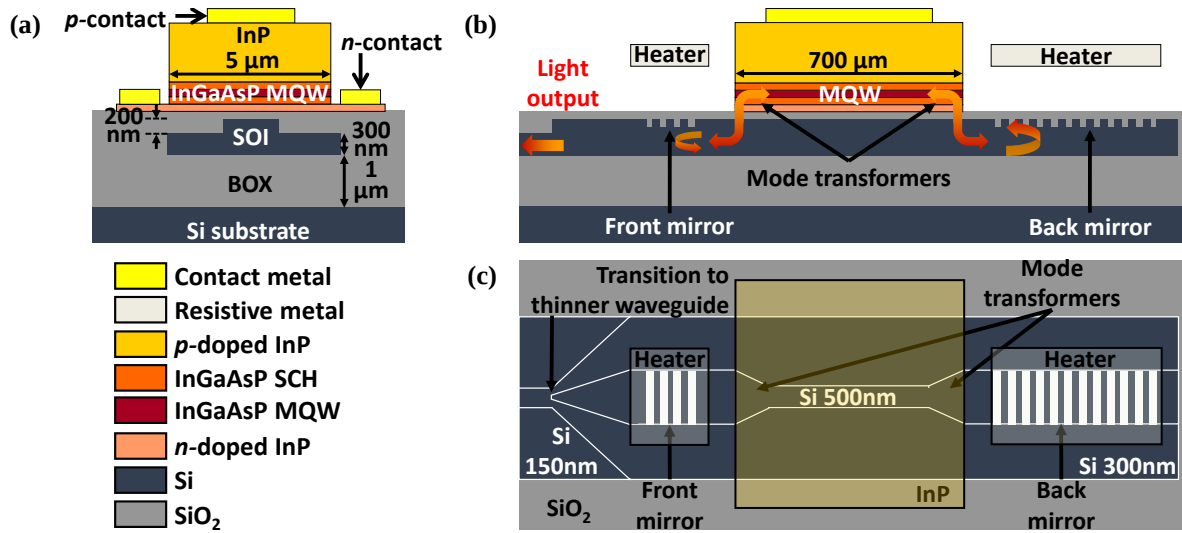


Figure III-1: (a) Transversal, (b) longitudinal, and (c) top schematic views of the laser.

The III-V active waveguide is formed by etching the bonded III-V stack down to the *n*-doped InP layer. As shown on **Figure III-1(a)**, the *n*-contact is deposited on the sides of the waveguide, thus the *n*-doped InP layer thickness can be reduced to facilitate the coupling between the two waveguides. On the opposite, the *p*-contact is

ⁱ Only the layers useful for the design of the laser are presented here, while the few missing ones are useful from a fabrication perspective. The complete III-V stack used for the demonstration, with the dimensions of each layer, is detailed in **section III-3.2**, on **Table II-5**.

located above the active gain region. Therefore, the p -doped InP layer must be thicker than $1.5\mu\text{m}$ ($2\mu\text{m}$ in our case), in order to separate the metal layer and the gain region, and to avoid large optical losses. This thick cladding layer is the main reason why the hybrid laser sources induce a large topography on the wafer.

The III-V active waveguide is $700\text{-}\mu\text{m}$ -long and $5\text{-}\mu\text{m}$ -wide. While it would have been preferable to test different gain lengths, this was impossible due to the large footprint of a single transmitter, which limited the maximum number of components on each die. The width is fixed as $5\mu\text{m}$, which facilitates the fabrication of the waveguides and the p -contact metallization, using the fabrication steps later shown in **section III-3.2**. Additionally, even if a narrower waveguide might reduce the threshold current, the mode would also become more sensitive to the scattering losses caused by waveguide sidewall roughness after etching, thus leading to the opposite effect.

Silicon waveguide description

For the hybrid III-V on silicon laser, adiabatic coupling is used rather than evanescent coupling, in order to avoid a trade-off between modal gain and light extraction from the active region. As explained in **section I-3.2**, in order to realize this coupling, the propagation constants ratio between the isolated III-V ($\tilde{\beta}_{III-V}$) and silicon ($\tilde{\beta}_{Si}$) waveguides must be inverted along the light propagation direction. The condition can also be expressed with the effective indices of the isolated III-V ($\tilde{n}_{III-V}$) and silicon ($\tilde{n}_{Si}$) waveguides. Therefore, below the gain region, the condition $\tilde{n}_{III-V} > \tilde{n}_{Si}$ must be respected to have maximum light confinement in the MQW, and at both ends of the active region, the condition must be switched to $\tilde{n}_{III-V} < \tilde{n}_{Si}$, to ensure light coupling into the silicon waveguide. These conditions can be achieved by reducing/increasing the width of the silicon waveguide, III-V waveguide or even both waveguides. However, the effective index of a III-V waveguide with dimensions above $1\mu\text{m}$ is quite high (typically above 3.2 [70], [144]). Thus, if the silicon waveguide is too thin, the conditions for coupling become difficult to reach. This can be an issue since components found in most silicon photonics platforms exhibit typical thicknesses below 300nm [11]–[14], as the silicon modulator presented in the previous chapter.

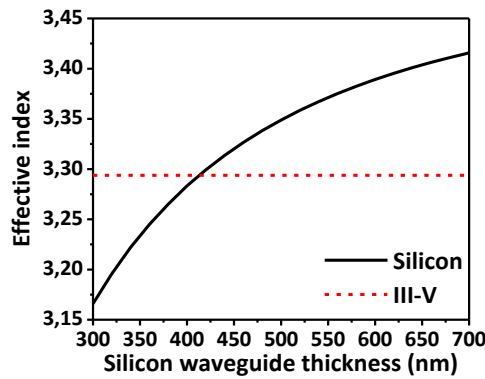


Figure III-2: Effective index of the simulated fundamental TE modes of the III-V waveguide and a silicon slab waveguide ($10\text{-}\mu\text{m}$ -wide) versus its thickness, at $1.31\mu\text{m}$. Both waveguides are cladded in SiO_2 .

Figure III-2 shows the effective index of the simulated fundamental TE mode for a $10\text{-}\mu\text{m}$ -wide silicon slab waveguide with different thicknesses, as well as the simulated effective index of the fundamental TE mode for our III-V waveguide ($\tilde{n}_{III-V} \approx 3.29$). Therefore, for a fixed silicon thickness, this will be the maximum effective index reachable for the silicon waveguide. It can be seen that for 300-nm -thick silicon waveguides, the effective index is limited to 3.17, which is too low. Even if it might be possible to reduce $\tilde{n}_{III-V}$ by narrowing the III-V waveguide, this solution also has several technological limitations, which are discussed later in **section IV-1.1**. Therefore, it is preferable to transfer a maximum of optical functions in the SOI layer. In our design, only the variation of the silicon waveguide width is considered, while the III-V waveguide width remains constant. Finally, 500-nm -thick silicon rib waveguides with a 300-nm -thick slab section are used for the lasers. This way, a large difference (positive or negative) between the effective indices of both waveguides can be reached to fulfil the conditions for adiabatic coupling, and the other silicon components can be fabricated in the 300-nm -thick slab section.

Passive components forming the laser cavity

As it can be seen in **Figure III-1(b)-(c)**, light is transferred by using tapers (also referred as mode transformers) in the rib section of the silicon waveguide, located at both ends of the III-V waveguide. The rib is narrowed down to $0.6\mu\text{m}$ below the III-V active region to ensure maximum light confinement in the MQW, but its width increases up to $1.55\mu\text{m}$ at both ends of the III-V mesa in order to couple light into the silicon waveguide. The transition is realized over $100\mu\text{m}$, using specific mode transformers, which shape is a bit more complex than the linear shape depicted on **Figure III-1(c)**. The design of these mode transformers is detailed in **section III-1.2**.

The device architecture is that of a DBR laser, where silicon Bragg mirrors are located at each end of the III-V active waveguide. In order to provide both high confinement and light extraction in a single direction, the back mirror should exhibit a reflectance as high as possible, while the front mirror should have a lower reflectance. These mirrors should also be selective in terms of wavelength to ensure single-mode emission. The main advantage of the DBR structure is its wavelength tuning capabilities. Indeed, by depositing a resistive metal layer above the gratings, and using it as an electrical heater, it is possible to increase the selected wavelength. The design of these mirrors is described in **section III-1.3**. Finally, once light is out of the laser cavity, it is transferred from a 500-nm -thick rib waveguide with a 300-nm -thick slab section into a 300-nm -thick rib waveguide with a 150-nm -thick slab section, as depicted on **Figure III-1(c)**. This transition is done so the laser can be used with other silicon photonics components designed on a 300-nm -thick SOI platform (such as the modulator presented in **chapter II**), and is detailed in **section III-2.1**.

III-1.2. Adiabatic taper design

In **section I-3.2**, we have seen that transferring the optical power from the III-V to the silicon waveguide using adiabatic coupling consists in shaping the even supermode from the state $[\delta < 0, |\delta| \gg \kappa]$ to the state $[\delta > 0, |\delta| \gg \kappa]$, where:

$$\delta = \frac{\tilde{\beta}_{Si} - \tilde{\beta}_{III-V}}{2} = \frac{\pi}{\lambda} (\tilde{n}_{Si} - \tilde{n}_{III-V}) \quad (\text{III.1})$$

To limit propagation losses, the mode transforming region should be as short as possible, while limiting the fraction of power scattered in other supermodes, referred as ε . To realize the coupling section as short as possible, Sun et al. proposed an adiabaticity criterion^j for the design of the taper [79], [80], which is:

$$\gamma = \frac{\delta(z)}{\kappa} = \tan \left[\arcsin \left(2\kappa\varepsilon^{1/2} [z - z_0] \right) \right] \quad (\text{III.2})$$

Where γ is a normalization parameter, κ is the coupling constant between the III-V and silicon waveguides, z is the position on the axis, and z_0 correspond to the phase-matching position, where $\gamma = 0$. Using **Eq. (III.2)**, it can be seen that if $(z - z_0)$ varies from $-\frac{1}{2\kappa\varepsilon^{1/2}}$ to $\frac{1}{2\kappa\varepsilon^{1/2}}$, γ (or δ) varies from $-\infty$ to $+\infty$, which means that the power in the even supermode is effectively transferred from the III-V to the silicon waveguide. Thus, a taper respecting this criterion and with a maximum length $L_{ac} = \frac{1}{\kappa\varepsilon^{1/2}} = 2z_0$ will transfer the mode with a fixed fraction of power ε^k in the odd mode. This length is necessarily larger than the length L_{dc} defined for directional coupling in **section I-3.2**, by a factor $\frac{\pi}{2\varepsilon^{1/2}}$.

Therefore, if the width of the silicon waveguide (W_{Si}) – which refers to the width of the rib section – can be expressed as a function of the parameter γ , the shape of a taper respecting the adiabaticity criterion can be found. To do so, we must find $\delta(W_{Si})$ and κ . The first can be found from the effective index of the eigenmodes for the isolated III-V and silicon waveguides, which can be calculated using a FEM mode solver, and are displayed on **Figure III-3(a)**. Only the first (referred as mode 0, or fundamental mode) and second TE modes of the silicon waveguide are displayed. Given its geometry, the TE modes of the III-V waveguide are close to each other. According to the coupling theory, only the modes with a sufficient overlap can couple with each other and form

^j In this criterion, the coupling constant κ is assumed to be constant along z .

^k It can also be noted that ε is different from the coupling efficiency of the taper.

supermodes for the hybrid waveguide structure. Thus, the even/odd TE modes of the III-V waveguide can only couple with the even/odd TE modes of the silicon waveguide. In our case, we only need to focus on the fundamental modes of each waveguide. The $\delta(W_{si})$ parameter between these two modes is calculated using Eq. (III.1), and presented on Figure III-3(b), with the phase-matching condition verified at $W_{si} = 0.85\mu m$.

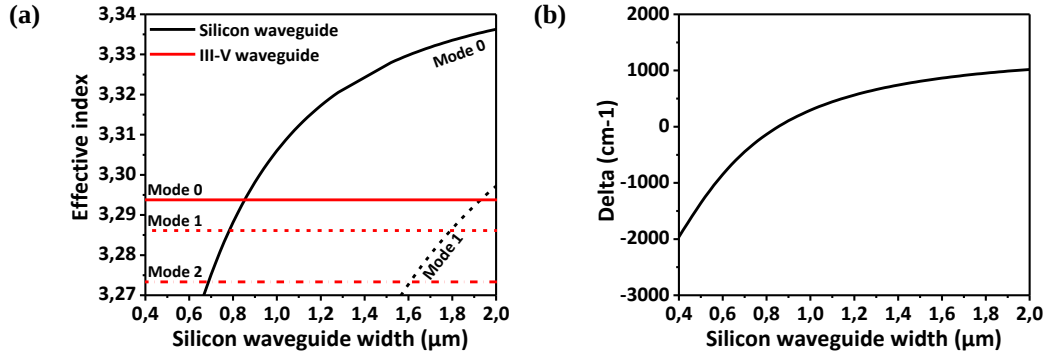


Figure III-3: (a) Effective index of the simulated TE modes of the isolated III-V and silicon waveguide versus the width of the silicon rib section, at $1.31\mu m$. Both waveguides are cladded in SiO_2 . (b) Propagation constant difference between the fundamental modes of each waveguide.

The coupling constant (κ) depends mainly on the gap between the two waveguides, can be found from the effective index of the even ($\tilde{n}_e$) and odd ($\tilde{n}_o$) supermodes for the hybrid III-V and silicon waveguide, which are calculated using the same FEM mode solver, and are displayed on Figure III-4(a). Using Eq. (I.19), we can see that at the phase matching point:

$$\kappa = \frac{\tilde{\beta}_e - \tilde{\beta}_o}{2} = \frac{\pi}{\lambda} (\tilde{n}_e - \tilde{n}_o) \quad \text{at } \delta=0 \quad (III.3)$$

Thus, we got $\kappa = 201.45 cm^{-1}$ for our structure with a 100-nm-thick oxide gap. The effective index of the TE modes for the isolated III-V and silicon waveguides are also represented Figure III-4(a). It can be seen that the even and odd supermodes follow the even TE modes of the isolated waveguides. However, the even supermode does not fit exactly the fundamental mode of the silicon waveguide after phase matching. This is the limit of the weak coupling theory, which shows that the coupling is actually quite strong. Thus, the value of κ is actually larger than expected. However this underestimation of κ is not an issue for the design of the taper. Figure III-4(b) shows the confinement factor of the even and odd mode as a function of W_{si} . It can be seen that for our structure, the confinement factor in the MQW (without the barriers) is expected to be 16.5% at most. When W_{si} increases, so does the even mode confinement in the silicon waveguide, while the odd mode stays mainly in the III-V waveguide. Since the DBR cavity is formed by reflectors in the silicon waveguide, the odd mode will need a much larger injected current than the even mode to reach the threshold condition. Most of the even supermode transformation happens for a silicon rib waveguide width between 0.6 and $1.2\mu m$.

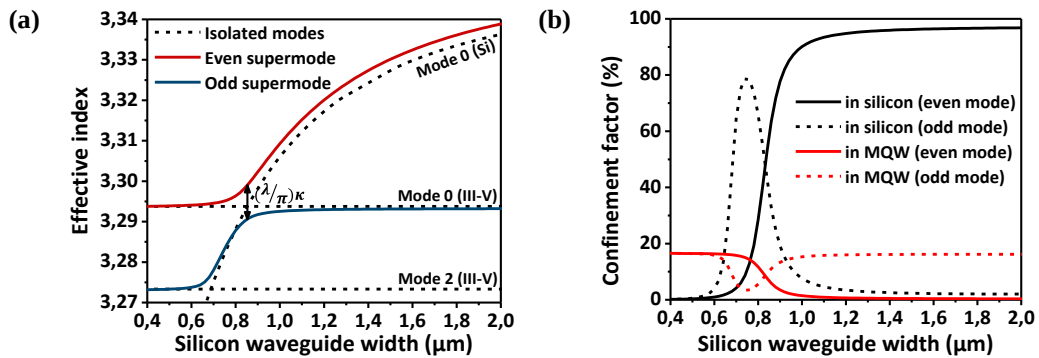


Figure III-4: (a) Effective index of the simulated even and odd supermodes of the hybrid III-V and silicon waveguide, and (b) confinement factor of each supermode in the silicon waveguide and in the MQW (without the barriers), versus the width of the silicon waveguide, at $1.31\mu m$.

Once $\delta(W_{si})$ and κ are known, the dependence of W_{si} on γ can be evaluated and is displayed on **Figure III-5**. By applying a fitting function, we can find W_{si} for any value of γ . The next step is to use **Eq. (III.2)** to have the function $\gamma(z)$ for an arbitrarily chosen value of ε . Given the value of κ at phase matching, it would require large ε values to have short coupling lengths. However, since κ is actually larger than its value at phase matching, ε will actually be lower than its chosen value.

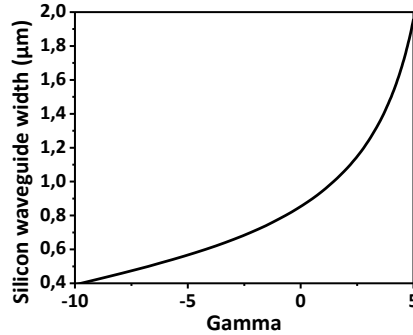


Figure III-5: Dependence of the silicon waveguide width on the γ parameter.

Once the function $W_{si}(z)$ is found for a fixed value ε , several taper shape respecting the adiabaticity criterion can be found by selecting a couple of input and output width for the silicon rib width. As it was shown on **Figure III-4(b)**, the even supermode is mainly confined in the MQW when $W_{si} < 0.6\mu\text{m}$. Therefore, it is preferable to use input widths close to this value. However, we will see later in **section IV-2.2** that tapers with a larger input width might be necessary in certain cases. The taper shape used for the DBR laser can be seen in **Figure III-6(a)**, where its width and length have been normalized for the next design step.

Once the shape of the taper is fixed, the next step is to optimize the input and output widths for a given length. In our case, a length of $100\mu\text{m}$ is targeted for the taper. To do so, the beam propagation method (BPM) is used to simulate the propagation of the even supermode along the taper. The percentage of optical power guided by the silicon waveguide at the taper output gives the coupling efficiency of the taper. The BPM calculation is realized for several couples of input and output widths. The taper dimensions are obtained by using the following formula:

$$W_{si}(z) = W_{input} + (W_{output} - W_{input})f(z) \quad (\text{III.4})$$

Where $f(z)$ is the normalized shape displayed on **Figure III-6(a)**. The result of the calculation is displayed on **Figure III-6(b)**. Several combinations of widths offer a coupling efficiency above 90%, with similar robustness. Finally, the input width is fixed at $0.6\mu\text{m}$, and the output width at $1.55\mu\text{m}$. While the output width is more robust than the input width towards lithographic and etching variations, both remains within a reasonable process variations window. Thus, for less than $\pm 30\text{nm}$ variation on the input width, the coupling efficiency remains above 90%.

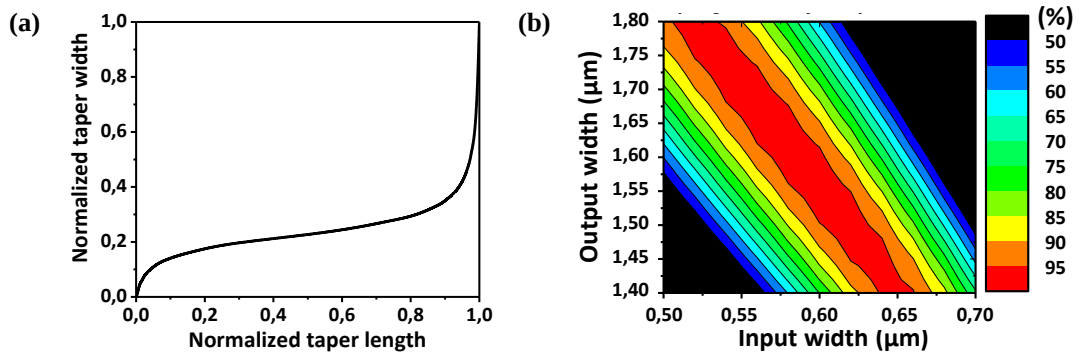


Figure III-6: (a) Normalized shape of the adiabatic taper. (b) Power confinement factor in the silicon waveguide at the end of the III-V region versus the input and output widths of the adiabatic taper.

The performances of this taper are displayed on **Figure III-7**. **Figure III-7(a)** shows the electrical field distribution of the fundamental even mode at the center of the III-V region, where the silicon waveguide is narrowest, and the mode overlaps mostly with the MQW. On **Figure III-7(b)**, the displayed field distribution corresponds to the end of the III-V region, where the silicon waveguide is widest, and the mode overlaps mostly with the silicon waveguide. **Figure III-7(c)** illustrates the power transfer between III-V and silicon waveguides along the propagation direction, while **Figure III-7(d)** shows the confinement factor in the silicon waveguide and MQW across the entire III-V region. Light confinement in the MQW is simulated to be $\approx 16.4\%$ at the center, while over 90% of the light is confined in the silicon waveguide at the end of the mode transformers.

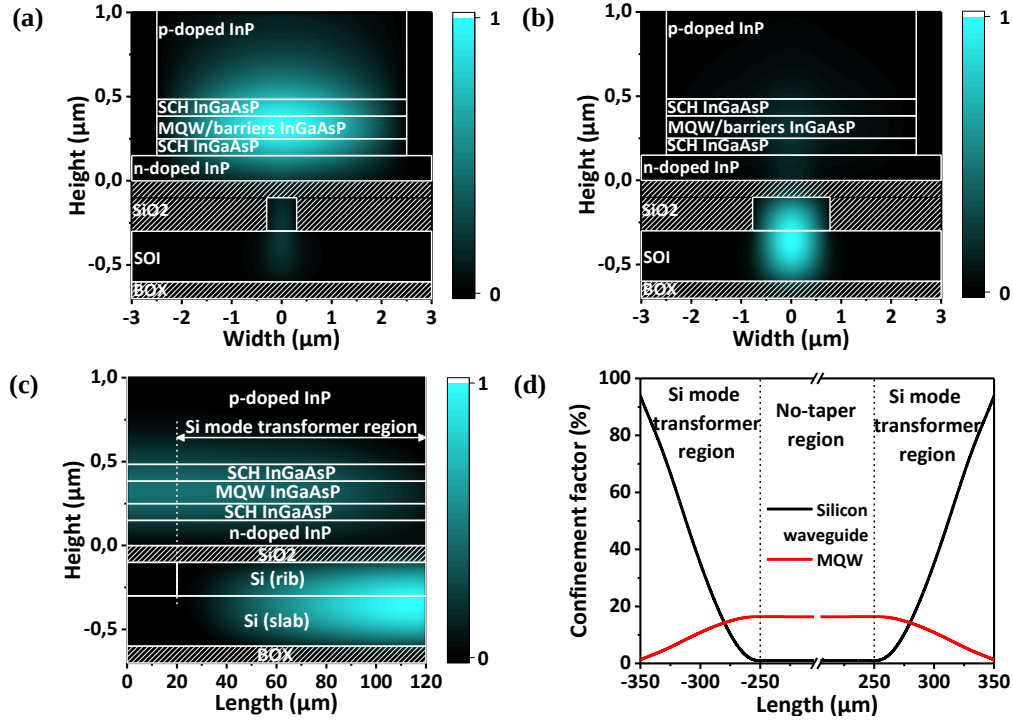


Figure III-7: Electric field distribution of the fundamental even mode at the center (a), and end (b) of the gain region. (c) Longitudinal view showing the power transfer from the end of the III-V region to the Si waveguide. (d) Confinement factor in the MQW and Si waveguide across the whole III-V region.

The robustness of the mode transformer is also evaluated by simulating the power confined in the silicon waveguide at the end of the III-V region with other geometrical variations. The influence of the mode transformer length is displayed on **Figure III-8(a)**. It can be seen that most of the optical power is transferred in the silicon waveguide at the targeted length of $100\mu\text{m}$, with reduced variations after this length. The influence of the SiO₂ gap between III-V and silicon is shown on **Figure III-8(b)**. This parameter is the most critical one, since for a $\pm 30\text{nm}$ variation from the targeted 100nm gap, the coupling efficiency falls under 70%.

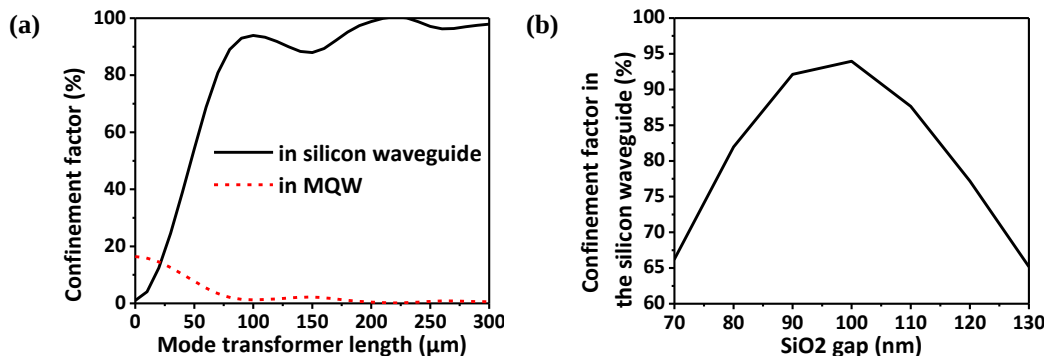


Figure III-8: Power confinement factor in the silicon waveguide at the end of the III-V region versus (a) mode transformer length, and (b) SiO₂ gap between III-V and silicon.

III-1.3. Reflectors design

This section is dedicated to the design of the reflectors used for the hybrid III-V on silicon DBR lasers. These reflectors are based on two Bragg gratings, located in the silicon waveguide, at each side of the III-V region, as shown on **Figure III-1**. Two 60- μm -long linear tapers are used to link the output of the 100- μm -long adiabatic tapers presented in **section III-1.2**, and the reflectors. Contrary to cleaved facets, the reflectance spectrum (defined as the modulus squared of the mirror reflection coefficient) of gratings is much narrower, which make them suited for single-wavelength operation.

In order to provide both high confinement and light extraction in a single direction, the back mirror should exhibit a reflectance as high as possible, while the front mirror should have a lower reflectance. These mirrors also act as filters, and have to be sufficiently selective in terms of wavelength to ensure single-mode emission. First the geometrical parameters and basic equations useful for the conception of the reflectors are briefly presented¹, followed by the design of the gratings used for the DBR lasers.

Geometrical parameters and basic equations

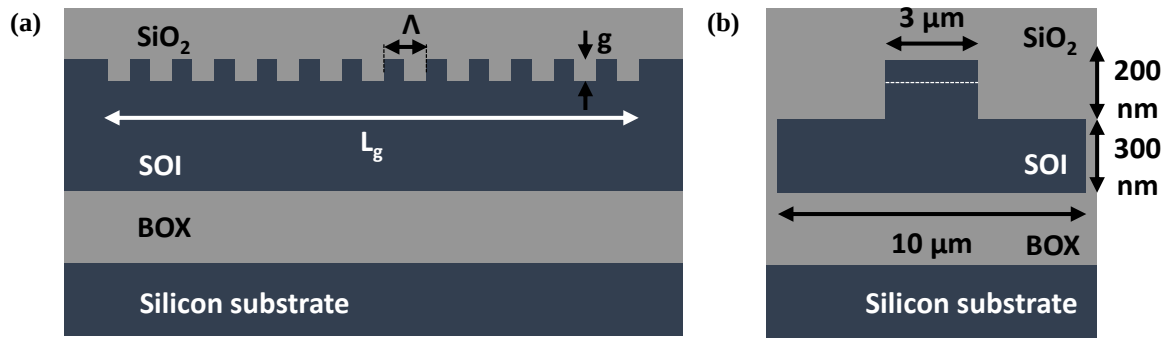


Figure III-9: Longitudinal (a) and transversal (b) schematic views of the Bragg reflectors.

Schematic views of the gratings are shown on **Figure III-9**. They are patterned in the previously described rib waveguide, with a fixed rib width of 3 μm . Only first order gratings, with a 0.5 duty cycle are considered. Thus, the gratings are essentially defined by their period (Λ), their etching depth (g), and their length (L_g). The influence of these parameters on the gratings performances (reflectance and selectivity) can be simply described by using a set of equations derived from the coupled-mode theory. First, the reflectance spectrum (R) of a Bragg grating can be expressed by:

$$R = \frac{\kappa_g^2 \tanh^2(s_g L_g)}{s_g^2 + \delta_g^2 \tanh^2(s_g L_g)} \quad (\text{III.5})$$

Where κ_g , s_g , and δ_g are respectively referred as the grating coupling constant, the decay constant and the detuning parameter from the Bragg condition. These parameters are defined by:

$$\kappa_g \approx \frac{2\Delta\tilde{n}}{\lambda_g} \quad (\text{III.6})^m$$

$$\delta_g = 2\pi\bar{n} \left(\frac{1}{\lambda} - \frac{1}{\lambda_g} \right) \quad (\text{III.7})$$

$$s_g = \sqrt{\kappa_g^2 - \delta_g^2} \quad (\text{III.8})$$

Where $\Delta\tilde{n}$ and $\bar{n}$ are respectively the difference and the average value of the effective index of the non-etched and etched cross-sections of the grating, and λ_g is the Bragg wavelength defined by:

¹ A complete description of the Bragg gratings based on the transmission matrix formalism and coupled-mode theory can be found in [3].

^m This approximation is only valid for gratings with a small index difference and a 0.5 duty cycle.

$$\lambda_g = 2\bar{n}\Lambda \quad (\text{III.9})$$

The reflectance spectrum is centered on this Bragg wavelength. Using **Eq. (III.5)** and **(III.7)**, it can be seen that when $\lambda = \lambda_g$ (Bragg condition), $\delta_g = 0$, and the reflectance reaches its peak value (R_p), defined as:

$$R_p = \tanh^2(\kappa_g L_g) \quad (\text{III.10})$$

It can be seen from **Eq. (III.10)** that the peak reflectivity increases with both L_g and κ_g , which – according to **Eq. (III.6)** – is proportional to the effective index difference $\Delta\tilde{n}$. $\Delta\tilde{n}$ increases with the etching depth, thus so does κ_g and R_p .

Etching depth

Examples of reflectance spectra for two gratings with different etching depths (shallowly and deeply-etched) are displayed on **Figure III-10**. Their respective lengths have been adjusted so their $\kappa_g L_g$ product is approximately the same. While having the same peak reflectance, the optical bandwidth of the deeply-etched grating is larger. This optical bandwidth is defined by the full width at half maximum (FWHM) of the reflectance spectrum, and shown on **Figure III-10**.

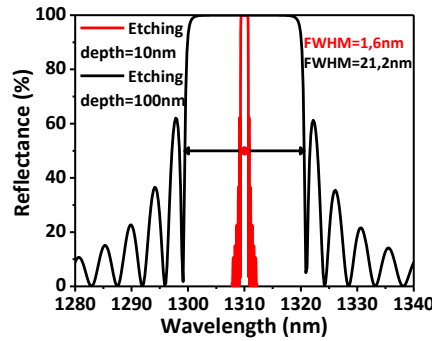


Figure III-10: Reflectance spectra for two gratings, with different etching depths, but the same product $\kappa_g L_g \approx 5.9$, thus the same peak reflectance. The targeted Bragg wavelength is 1310nm for each grating.

For $g = 10\text{nm}$ (red line), $\kappa_g \approx 79\text{cm}^{-1}$, and $L_g = 750\mu\text{m}$.

For $g = 100\text{nm}$ (black line), $\kappa_g \approx 1170\text{cm}^{-1}$, and $L_g = 50\mu\text{m}$.

If we do not take into account the effective length added by the gratings – defined later in this section, and which reduces even further the FSR –, we see that in our case the FSR will be approximately below 0.26nm using **Eq. (I.9)** (calculated with $\lambda = 1.31\mu\text{m}$, $n_g \approx 4$, $L_a = 700\mu\text{m}$, $L_p = 120\mu\text{m}$). Thus, if used as reflectors and as wavelength selection filters, deeply-etched mirrors such as the one displayed on **Figure III-10** are not suited for single-wavelength operation, but for Fabry-Pérot cavities (such as in [46]). Therefore, the rest of the study is based on shallowly-etched silicon gratings.

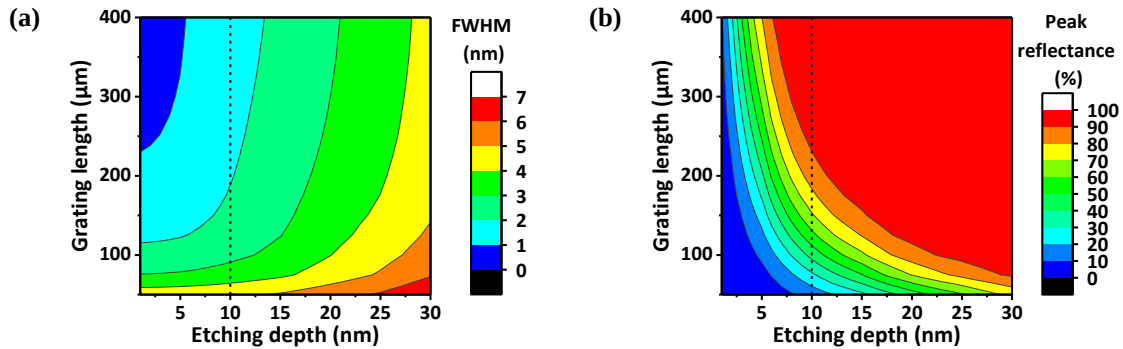


Figure III-11: (a) FWHM and (b) peak reflectance versus the etching depth and grating length of silicon Bragg gratings. The targeted Bragg wavelength is 1310nm for each grating.

Figure III-11(a) displays the *FWHMs* for several gratings with different etching depths and grating lengths. While the *FWHM* depends on both parameters, it can be seen that its main dependence is on the etching depth (thus on κ_g). While it might be tempting to reduce the etching depth to increase the selectivity of the grating, we must not forget that a grating with a large reflectivity (as close as possible to 100%) is needed for the laser. Reducing the etching depth also results in longer gratings necessary to reach the desired peak reflectance, as shown on **Figure III-11(b)**. Longer gratings will increase the already large footprint of the DBR laser, and increase the risk of a defect appearing during the fabrication of the grating. Finally, a trade-off etching depth of 10nm is used for the gratings, which corresponds to a coupling constant of 79cm^{-1} .

Grating period

The influence of the grating period on the Bragg wavelength is depicted on **Figure III-12(a)**. A linear dependency is obtained, with grating periods slightly below 196nm for a 1310nm operation. It can be seen that a $\pm 1\text{nm}$ variation on the grating period results in a $\pm 6.2\text{nm}$ shift on the Bragg wavelength. Since the wavelength spacing targeted for the transmitter is $\approx 4.5\text{nm}$, it will not be possible to achieve it by simply using different grating periods. However, it will become possible by using thermal tuning. By injecting a current in the heaters present above the reflectors and depicted on **Figure III-1**, we will be able to locally increase the effective index of refraction of the grating, thanks to the thermo-optical effect described in **section II-1.3**, and thus to increase the Bragg wavelength [see **Eq. (III.9)**]. While variations on the grating period cannot be used to fit all the wavelengths targeted by the transmitter, it can be used for a coarser tuning of the wavelength. In order to demonstrate transmitters with different wavelengths, two periods of 195nm and 197nm are selected, giving Bragg wavelengths of 1303.3nm and 1315.7nm respectively.

It can be noticed that according to the same equation, variations on the etching depth also have an influence on the Bragg wavelength, as shown on **Figure III-12(b)**. However, this influence is quite weak, since a $\pm 1\text{nm}$ variation on the grating period results in a $\pm 0.1\text{nm}$ shift on the Bragg wavelength.

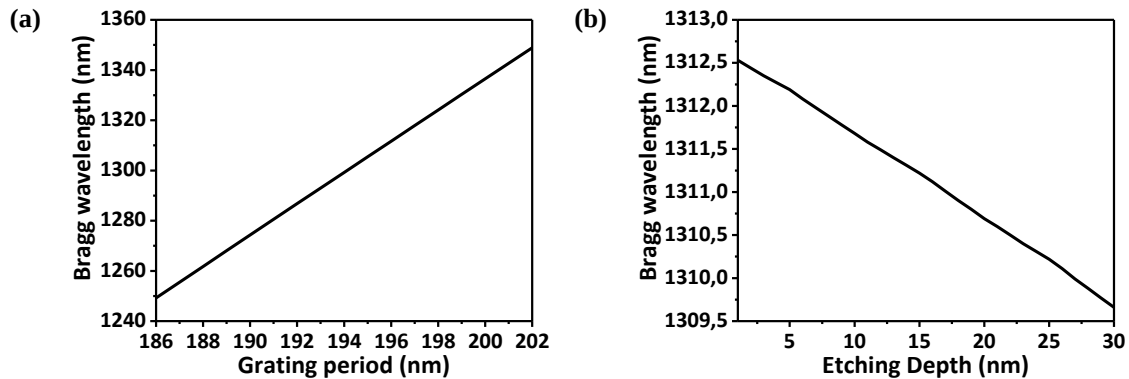


Figure III-12: Bragg wavelength dependency on (a) grating period for a 10nm etching depth, and (b) etching depth for a 196nm grating period.

Grating length

Once the etching depth and grating period are fixed, the last step is to fix the lengths of the front and back mirrors. However, variations on the etching depth will have a large effect on the peak reflectance of the mirrors. Due to the period of the gratings (below 200nm), it will be extremely difficult to monitor precisely the etching depth in spacings narrower than 100nm .

Figure III-13(a) shows the effect of a $\pm 4\text{nm}$ variation on the etching depth on the grating reflectivity spectrum, at fixed grating length of $250\mu\text{m}$. It can be seen that in these conditions, the reflection varies from 65 to 99%. To increase the gratings robustness, the peak reflectivity evolution with the grating length was studied, for different etching depths around 10nm . The results are displayed on **Figure III-13(b)**. For the back mirror, the grating length is fixed at $700\mu\text{m}$ to ensure a near 100% reflectivity, in case the grating is not sufficiently etched. For the front mirror, the grating length is fixed at $150\mu\text{m}$, to obtain a reflectivity between 35% and 85% for $\pm 4\text{nm}$ variation on the etching depth. Finally, the reflectance spectra of the designed front and back mirrors are displayed on **Figure III-14** for each grating period.

III-2. Transmitter structure

Once the DBR lasers are designed, the next step is to assemble all the elements described up to now, in order to form the hybrid III-V on silicon transmitters. In this section, the silicon thickness transition between the laser and the other silicon components (modulator, grating coupler) is presented, followed by the expected modifications on the silicon modulator performances due to its co-integration with the laser, and the complete layout of the transmitter.

III-2.1. Silicon thickness transition

As explained in **section III-1.1**, to reach an efficient coupling between III-V and silicon, it was necessary to use 500-nm-thick silicon waveguides, which are thicker than the silicon waveguides used for the silicon components introduced in **chapter II**, or in most silicon photonics fabrication platforms [11]–[14]. Rather than re-designing all other components with 500-nm-thick silicon waveguides, a transition in the silicon is used, to pass from a 500-nm-thick rib waveguide with a 300-nm-thick slab section, to a 300-nm-thick rib waveguide with a 150-nm-thick slab section. This transition is displayed on **Figure III-15**.

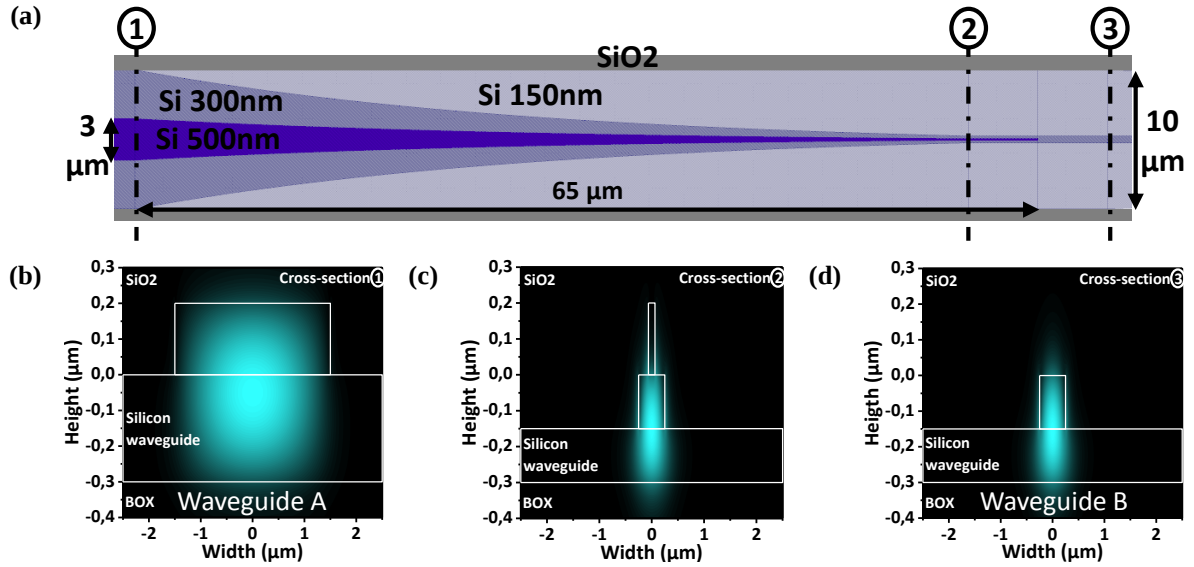


Figure III-15: (a) Top view of the transition. (b)-(d) Electric field distributions of the fundamental TE mode across the transition.

The thick-to-thin silicon transition is 65-μm-long, starts with a 3-μm-wide and 500-nm-thick rib waveguide (referred as waveguide A), and ends with a 500-nm-wide and 300-nm-thick rib waveguide (referred as waveguide B). At first, the width of the rib waveguide A is decreased using an exponential taper from 3 μm to 125 nm. In the same way, the top 150 nm of the slab section is reduced using the same taper shape from 10 μm to 500 nm, to form the rib of the waveguide B. After 65 μm, the 125-nm-wide and 200-nm-thick silicon tip stops. **Figure III-15(b)-(d)** shows electric field distributions of the fundamental TE mode across the transition. As it can be seen by comparing the cross-sections 2 and 3, the mode is well confined in the waveguide B, even with the 125-nm-wide and 200-nm-thick silicon tip. This confinement is evaluated at 90% in both cases, and the overall power loss due to the transition is expected to be as low as 1%.

While this solution will permit to use silicon photonics components designed for a 300-nm-thick waveguide with the hybrid III-V/silicon laser, it does not change the fact that a thicker SOI must be used. This remains an issue for components planned on thin SOI, since the additional etching processes would induce surface roughness and thickness variations, which in turn would reduce the devices optical performances and reliability. An alternative solution to this problem is presented later in **Chapter IV**.

III-2.2. Silicon modulator

The modulator integrated in the hybrid III-V on silicon transmitter has the same layout that the one presented in **section II-2.2**. However, due to the fabrication process of the transmitter – which is detailed later in **section III-3.2** –, the electrodes are slightly different: they are based on gold rather than AlCu, are thicker ($1.3\mu\text{m}$ instead of 650nm), and are also located further from the substrate due to a thick SiN cladding ($2\mu\text{m}$) deposited before the electrodes. As for the wafer used for the transmitter fabrication, they not only have a thicker SOI layer (500nm instead of 300nm), but also a thicker BOX layer ($1\mu\text{m}$ instead of 800nm). The modifications on the profile of the electrodes are illustrated on **Figure III-16**. Therefore, we performed additional RF simulations in order to verify the impact of these changes on the electrodes performances.

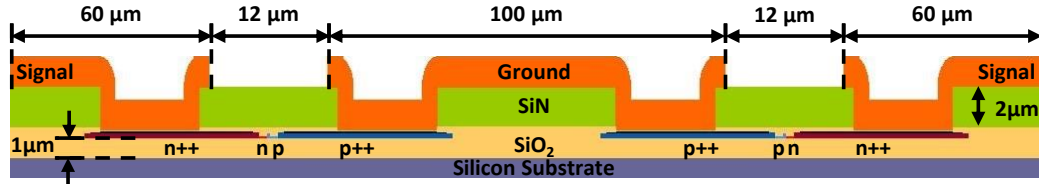


Figure III-16: Schematic transversal view of the silicon MZM, with its electrodes modified due to the co-integration with the hybrid III-V on silicon laser.

The R , L , G , and C parameters of the modified electrodes are compared to the previous ones presented in **section II-2.2** and displayed on **Figure III-17**. As it can be seen, the resistance and inductance of the line become smaller in the modified configuration, which is mainly due to the thicker metal layer used for the electrode. The conductance of the modified electrodes remains globally unchanged, while the capacitance is slightly increased due to a larger capacitance between the ground and signal lines. While not being visible on the graph, the larger distance between the substrate and the electrodes has also reduced the substrate capacitance, and thus the total capacitance at low frequency (below 1MHz).

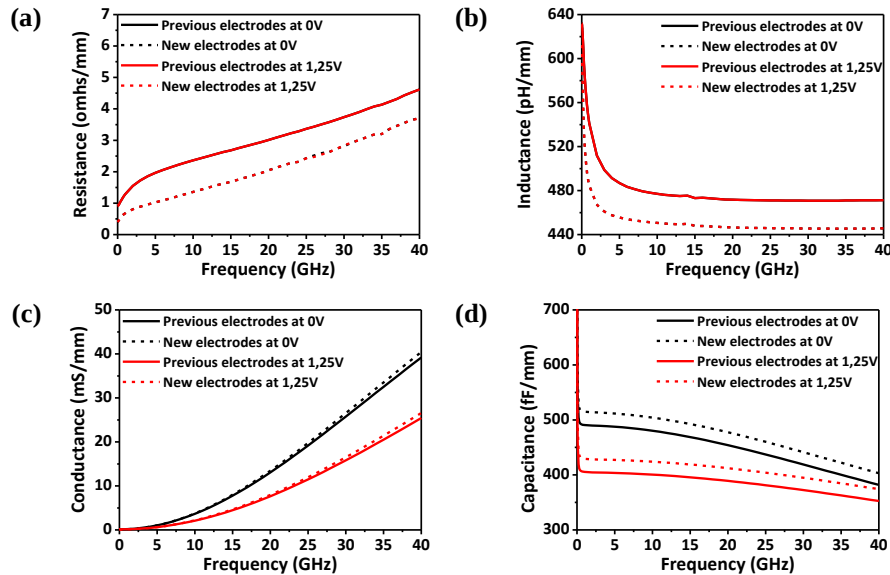


Figure III-17: Simulations of the linear (a) resistance and (b) inductance of the electrodes, and of the linear (c) conductance and (d) capacitance between the electrodes of the silicon modulator. All simulations are realized with a HR substrate ($750\Omega\cdot\text{cm}$). Lines: previous electrodes described in **section II-2.2**. Dashed lines: New electrodes due to the laser co-integration.

As in **section II-2.2**, the R , L , G , and C parameters can be used to directly evaluate the modulation depth of the line. The previous values of junction capacitances and series resistance used for the calculation are also used here. The results are displayed on **Figure III-18**, where it can be seen that the variations due to the electrodes fabrication process should have a negligible effect on the electro-optical bandwidth of the modulator. Therefore, it was not necessary to adapt the design of the modulator.

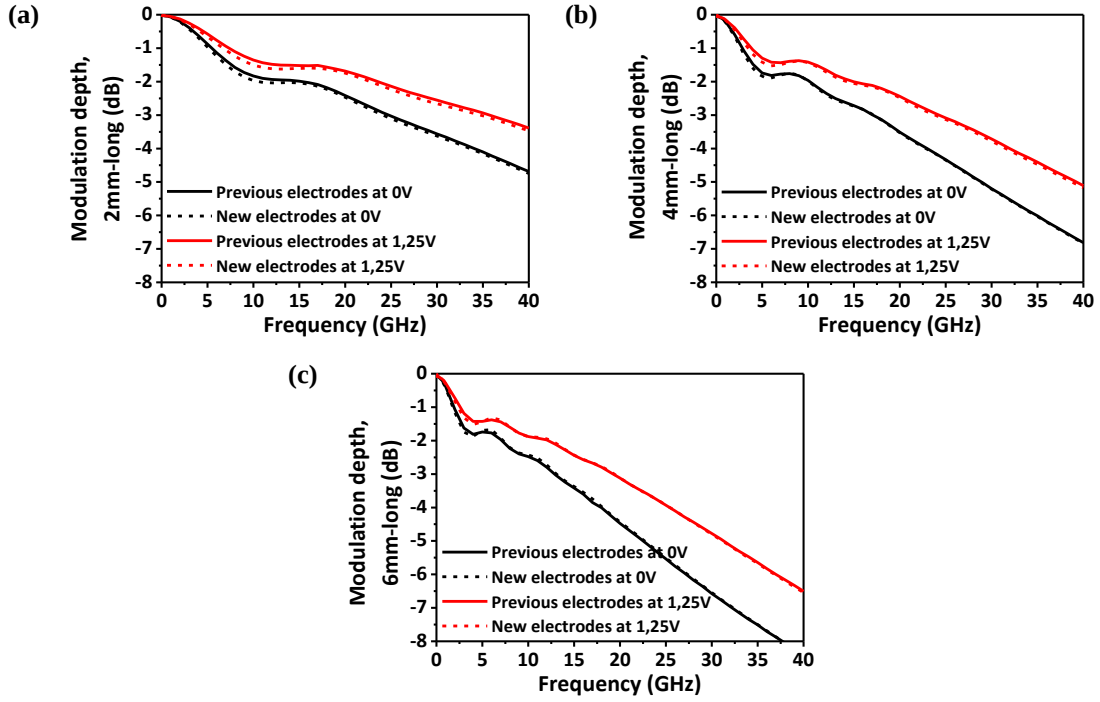


Figure III-18: Simulations of the modulation depth for the different lengths of the active region: (a) 2-*mm*-long, (b) 4-*mm*-long, and (c) 6-*mm*-long . All simulations are realized with a HR substrate ($750\Omega \cdot \text{cm}$). Lines: previous electrodes described in section II-2.2. Dashed lines: New electrodes due to the laser co-integration.

III-2.3. Transmitter layout

The layout of the complete hybrid III-V/silicon transmitter is shown on **Figure III-19**. The transmitter co-integrates the previously described hybrid III-V/silicon DBR laser and silicon MZM. Lasers with grating periods of 195 and 197nm have been assembled with 2, 4, and 6-*mm*-long MZMs, but only the 2 and 4-*mm*-long ones have been tested. Light emitted from the laser passes through the MZM, and is collected out of the wafer using a waveguide-to-fiber grating coupler whose peak wavelength is centered at 1.3 μm .

Moreover, two pairs of heaters are located above the Bragg mirrors and the 1-*mm*-long section of passive waveguide present in both MZM arms. The first pair of heaters is used to tune the laser emitting wavelength, while the second pair can set the phase difference between the MZM arms, to set it in a quadrature condition for any wavelength, as explained in section II-2.3.

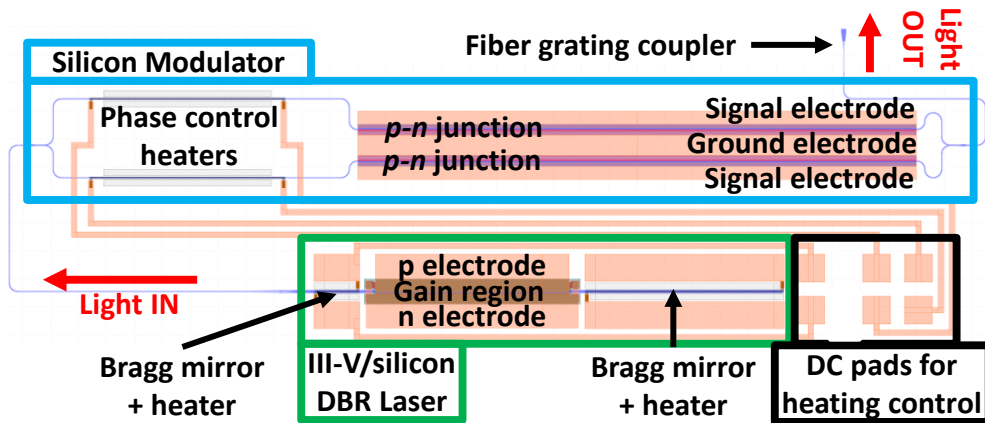


Figure III-19: Layout view of the complete transmitter co-integrating the hybrid III-V on silicon DBR laser and the silicon MZM (2-*mm*-long).

III-3. Hybrid III-V on silicon transmitter fabrication

The complete fabrication process of the hybrid III-V on silicon transmitter is illustrated on **Figure III-21**. These devices are fabricated on 8" SOI wafers from SOITEC, with a 500-nm-thick silicon layer, a 1- μm -thick BOX, on a HR handle silicon wafer (with a resistivity of $750\Omega\cdot\text{cm}$). A part of this fabrication process has been detailed in **section II-3**, and will not be re-detailed here. Therefore, this section is more focused on the steps related to the silicon components used for the lasers (Bragg reflectors, adiabatic tapers and transition from thick to thin silicon waveguides), and the steps following the III-V wafer bonding.

III-3.1. SOI part fabrication

Silicon patterning of the laser components [see **Figure III-21(a)-(b)**]

The fabrication starts with the patterning of the Bragg reflectors necessary to form the DBR laser cavity. These reflectors are realized through Electron Beam Lithography (EBL), and patterned in the silicon layer using a 10-nm-deep reactive ion etching. Given the period of the Bragg reflectors (195nm and 197nm), the gratings could have been patterned by 193nm DUV photolithography, but given the width of the grating (3 μm), preliminary tests showed that some of the lines formed during the photolithography of the resist layer collapsed along the grating. On the other hand, EBL uses a different resist, with a reduced thickness compared to 193nm DUV (200nm vs 400nm), which did not collapse after being patterned. The main issue with EBL is the risk of not repeating perfectly the period of the grating. Given the length of the mirrors (150 μm and 700 μm), this risk is not negligible, and can generate defects in the gratings.

Once the gratings have been etched and the resist stripped, a 160-nm-thick SiO₂ hard mask is deposited. Then, 193nm DUV photolithography, followed by hard mask etching, resist stripping, and a 200-nm-deep RIE are performed to define the rib waveguides used for the lasers. The adiabatic tapers for III-V to silicon coupling, as well as the transition from thick (500nm) to thin (300nm) silicon waveguides are defined during this step. SEM views of these two components, as well as the Bragg reflectors patterned with the previous etching step are displayed on **Figure III-20**.

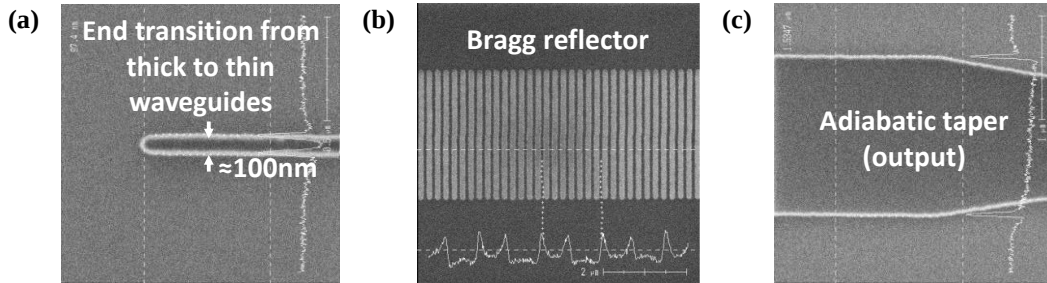


Figure III-20: SEM views of the (a) end transition from thick to thin silicon waveguides, (b) Bragg reflector, and (c) adiabatic taper used for the coupling between III-V and silicon waveguide.

Ion implantation, thin waveguide patterning, and silicide formation [see **Figure III-21(c)-(d)**]

After the patterning of the thick rib waveguides, a 5-nm-thick rapid thermal oxidation of the SOI layer is performed. From here, the fabrication steps described in **section II-3** are re-used, with the same ion implantation conditions to form the p - n junctions, the same etching conditions to pattern the thin rib waveguides used for the modulator and the other passive components, as well as the same annealing conditions for the dopants activation. The fabrication steps for the formation of the silicide layer on the modulator contact regions are also conserved.

SiO₂ encapsulation and planarization [see **Figure III-21(e)-(f)**]

While the encapsulation and planarization processes are technically the same that those used in **section II-3**, these steps were taken with greater care than for the stand-alone modulator. Indeed, as discussed in **section III-1.2**, the SiO₂ gap between the III-V and thick silicon waveguide must be as close as possible to its targeted value of 100nm, in order to keep a good coupling efficiency between the III-V and SOI waveguides. However, the thick encapsulation and planarization processes used for the fabrication of the transmitter both induce a non-uniformity on the SiO₂ gap between the center and the edge of the 8" SOI wafer.

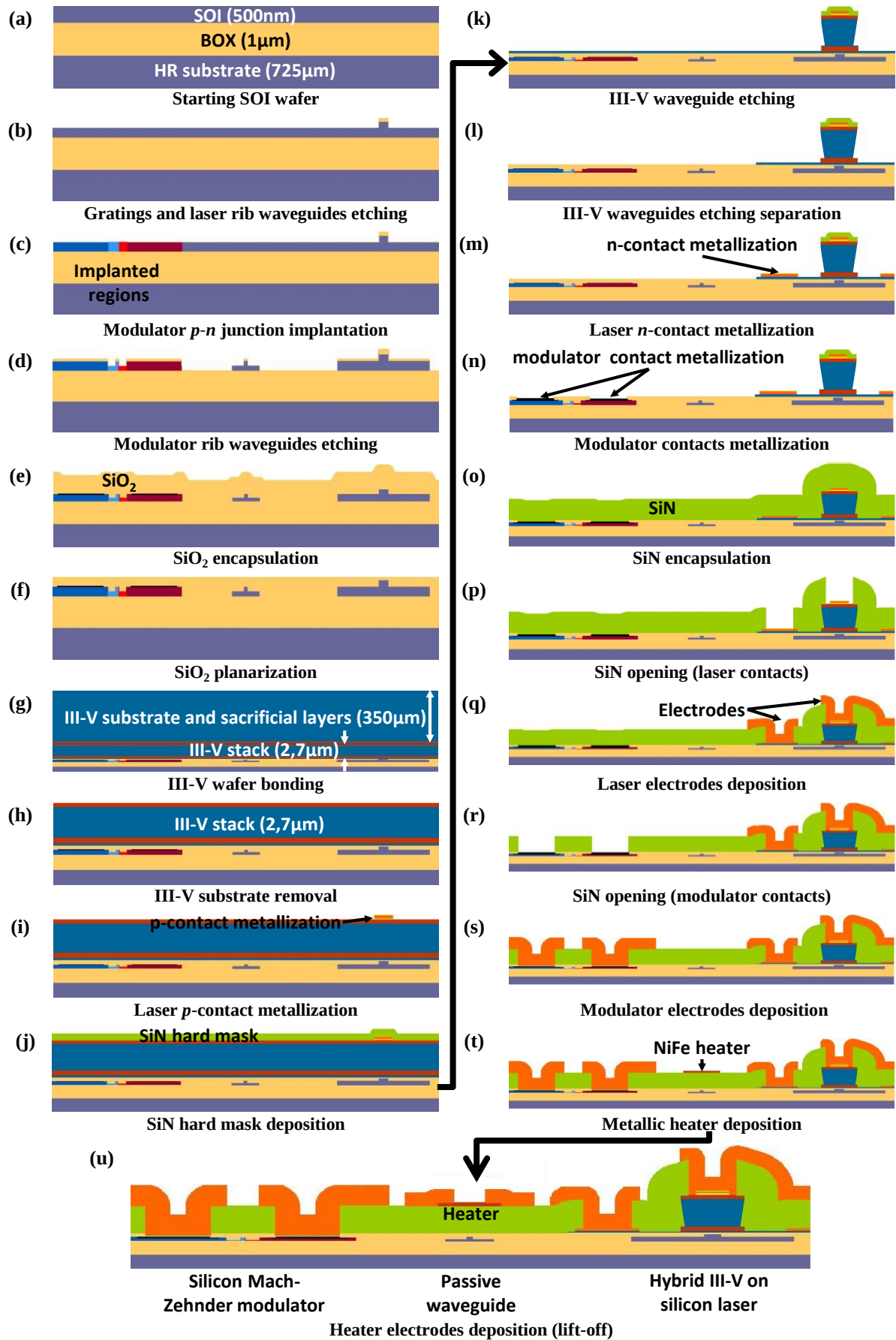


Figure III-21: Complete fabrication process flow of the hybrid III-V on silicon transmitters.

With the deposition process used for the encapsulation step, the SiO₂ thickness can vary up to $\approx 8\%$ of its targeted value between the center and the edge of the wafer. For instance, if the targeted thickness is $1\mu\text{m}$, this value is reached at the center, but approximately 920nm are deposited on the edges of the wafer. During the CMP step, the SiO₂ is also removed faster at the edges of the wafer than at its center, which increases further the non-uniformity. Therefore, the deposited SiO₂ thickness must be reduced to limit the non-uniformity on the wafer, but must stay thick enough to remove the topography over the wafer, with enough oxide remaining above the silicon waveguides. The topography left after the different etching steps and hard mask depositions was measured to be approximately 660nm . After several tests, the minimum deposited SiO₂ thickness necessary to remove this topography and leave more than 100nm over the thick silicon waveguides was estimated at 800nm . The planarization process was then maintained for a few seconds, until the desired SiO₂ thickness above the waveguide was reached. Using this method, the SiO₂ gap is estimated at 100nm at the center of the $8''$ SOI wafers, and down to 20nm at their edges. It is essential to end the process of the SOI wafers with a CMP step, in order to provide a clean surface for the III-V wafer bonding.

III-3.2. III-V integration

Wafer bonding and substrate removal [see Figure III-21(g)-(h)]

Once the SOI wafers have been patterned and planarized, the next step is to bond the III-V wafers. The composition of the III-V wafers used for the hybrid III-V on silicon transmitter is displayed on **Table III-2**. These wafers have been ordered to Landmark, a manufacturer of InP and GaAs-based epitaxial wafers. The epitaxial structure used for the laser is grown on top of a $2''$ InP wafer. The grey layers will be removed during the processing steps, while the other layers are used to form the III-V active waveguide. Amongst them, we can find the layers presented earlier in **section III-1.1**, such as the *n*- and *p*-doped InP layers, and the InGaAsP MQW and SCH.

Table III-2: III-V epitaxy layer structure used for the integrated III-V on silicon transmitter.
The grey layers are removed after the bonding, and only the other layers are used to form the active waveguide.

Layer	Material	Photoluminescence peak wavelength (μm)	Bandgap (eV)	Thickness (nm)	Doping (cm^{-3})
Substrate	InP			$350 (\mu\text{m})$	2×10^{18}
Transition	InP			50	Undoped
Stop-etch	InGaAs			300	Undoped
Sacrificial layer	InP			300	Undoped
<i>p</i> -doped contact	InGaAs	1.65	0.75	200	2×10^{19}
Transition	InGaAsP	1.1	1.13	50	5×10^{18}
<i>p</i> -doped cladding	InP	0.92	1.35	2000	$2 \times 10^{18} \rightarrow 5 \times 10^{17}$
SCH	InGaAsP	1.1	1.13	100	Undoped
Barriers (x7)	InGaAsP	1.1	1.13	10	Undoped
Wells (x8)	InGaAsP	1.29	0.96	8	Undoped
SCH	InGaAsP	1.1	1.13	100	Undoped
<i>n</i> -doped contact	InP	0.92	1.35	110	3×10^{18}
Super-lattice (x2)	InGaAsP	1.1	1.13	7.5	3×10^{18}
Super-lattice (x2)	InP	0.92	1.35	7.5	3×10^{18}
Bonding interface	InP	0.92	1.35	10	Undoped

It can be seen that a heavily *p*-doped InGaAs layer is also present above the *p*-doped cladding layer, and is used as the contact layer. Indeed, while a good contact resistance with metallic layers can be easily obtained on moderately *n*-doped InP (on the order of 1 to $5 \times 10^{18}\text{cm}^{-3}$) [145], [146], it is not the case for the *p*-doped InP layer. Thus, a highly *p*-doped InGaAs layer (on the order of $2 \times 10^{19}\text{cm}^{-3}$) is generally added to obtain a good contact resistance [146]–[148]. Since this highly doped layer would induce large optical losses through free-carrier absorption, it must be placed far from the gain region. This explains why the *n*-doped InP layer is placed between the MQW and the silicon waveguide in hybrid III-V on silicon structures. Super-lattice layers are also present, to stop the dislocations which might come from the bonding interface [149].

Both the III-V and SOI wafers surfaces are activated by an oxygen plasma, and then bonded at room temperature, followed by a 120 minutes post-bonding annealing at 300°C . After bonding, the InP substrate is removed using HCl/H₂O wet etching. The etch-stop InGaAs layer is used to control this etching step and protect the rest of the stack, since the HCl/H₂O solution is highly selective towards InGaAs. This etch-stop layer is then removed by H₂SO₄/H₂O₂/H₂O wet etching, and stopped by an InP sacrificial layer. This layer is normally intended to protect the p-doped InGaAs contact layer if an H⁺ implantation step is used to electrically isolate the edges of the III-V waveguide. The implantation step was not used for the transmitter, therefore this layer was directly removed by HCl/H₂O wet etching. This is not the case for the laser fabrication described in **section IV-3**, where this layer is fully used. Once this layer is removed, only the $\approx 2.7\text{-}\mu\text{m}$ -thick epitaxial structure used for the III-V waveguide remains, as shown detailed on **Table III-2**. The bonding defects are also revealed at this step, since only the successfully bonded areas will remain on the SOI wafer. To pursue the fabrication, the SOI wafers are then downsized from 8" to 3" around the bonded 2" III-V wafers, in order to enable subsequent processing in a standard III-V platform.

In **section III-3.1**, we saw that the SiO₂ gap between III-V and silicon waveguides was optimized at the center of the 8" SOI wafer. Originally, a single III-V wafer was supposed to be bonded at the center of each SOI wafer. Unfortunately, due to an incident during the fabrication, almost all of the SOI wafers processed for the transmitter were damaged, and only one of them remained for bonding. Since there was a large risk of losing again the processed wafer, rather than bonding a single III-V wafer at the center of the remaining SOI wafer, we decided to bond three III-V wafers on it, as shown on **Figure III-22(a)**. After the removal of the substrate and the sacrificial layers displayed on **Figure III-22(b)**, almost the whole surface of the III-V stack was left on each wafer, demonstrating an excellent bonding yield. While this method gave us three downsized wafers instead of one, the main issue is that the targeted SiO₂ gap between III-V and silicon waveguides is no longer respected. Given the location of the wafers, the average gap was estimated to $\approx 60\text{nm}$. Amongst the three wafers, only two have been through the complete process flow. One was used as a pathfinder for the other, which is shown on **Figure III-22(c)**, and which electro-optical performances are presented in **section III-4**.

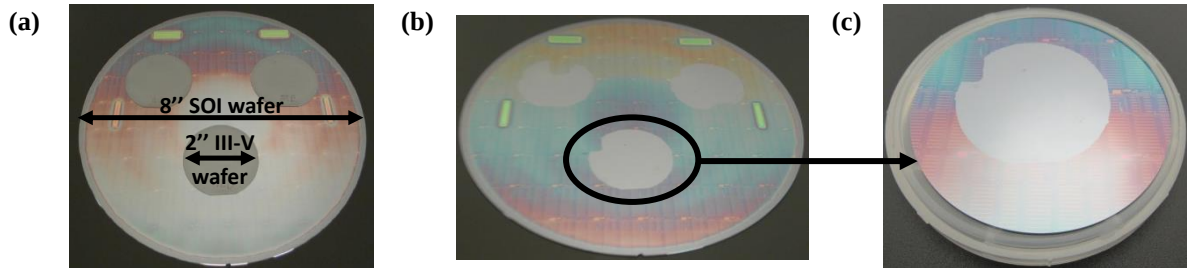


Figure III-22: (a) Top view of the 8" SOI wafer after III-V bonding, and (b) after the removal of the substrate and the sacrificial layers. (c) Top view of the wafer used for electro-optical characterization (presented in **section III-4**), after wafer downsizing.

Laser p-contact metallization [see **Figure III-21(i)**]

The *p*-contact is first formed by using the following metallic layers: Pd(10nm)/ Ti(30nm)/ Pt(80nm)/ Au(150nm)/ Pt(20nm). After a short annealing below 400°C , the Pd layer is expected to diffuse and intermix with the InGaAs layer, thus reducing the contact resistance, while the Ti layer improves the adhesion of the contact layers [148]. The first Pt layer is a barrier to prevent the diffusion of the Au layer used for contact reprimar in a later fabrication step [150]. The second Pt layer is present to protect the Au layer during the contact re-opening, also described later. The main interest of realizing the *p*-contact layer first is to have a clean InGaAs surface for the contact.

Given the composition of the layers, the best method to pattern the contacts is the lift-off process, which principle is reminded on **Figure III-23**. A bi-layer of lift-off resist (or LOR, not illustrated on **Figure III-23**) and negativeⁿ photoresist are first deposited and patterned. The LOR is easier to remove, making it suited for lift-off operation. As shown on **Figure III-23(a)**, the slope of the negative photoresist is inverted compared to positive resist, which limits the deposition of metal between the contact surface and the top of the resist, and facilitates

ⁿ In this document, if the polarity of the resist is not precised, it is automatically assumed to be positive.

the lift-off process. A wet cleaning process (HCl-based) is realized before the deposition, as well as an in-situ dry etching before the metal deposition, to remove any native oxide. The deposition is realized by evaporation, which deposit material from a single direction, and is thus suited for lift-off operation, since there will be no metallic connections between the contact surface and the top of the resist, as displayed on **Figure III-23(b)**. Finally, the resist bi-layer is removed by using an acetone solution for a few hours, and an ultrasonic bath if needed, leaving only the metallic layers on the contact region, as it can be seen on **Figure III-23(c)**.

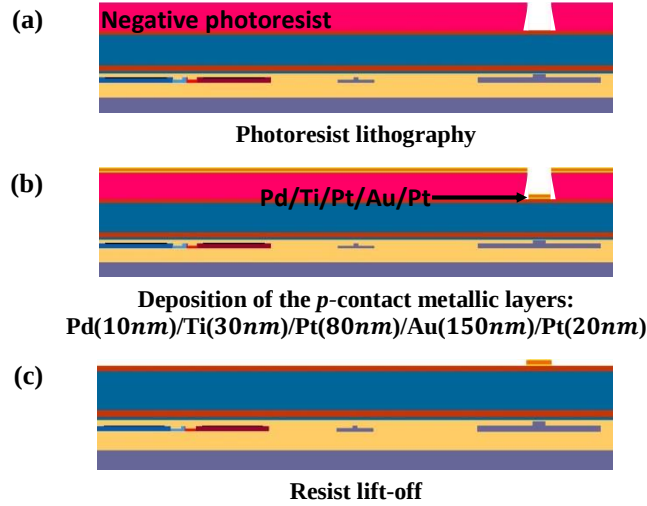


Figure III-23: Deposition and patterning of the *p*-contact metallic layers.

Top views of a *p*-metallic contact after deposition are shown on **Figure III-24**. No defects can be seen after the lift-off process. The annealing of the metallic contact layer is not realized directly after the deposition, since there is a potential risk of to peel-off the bonded III-V stack at annealing temperatures above 400°C [151]. Therefore, the contact annealing is realized after the patterning of the III-V waveguides.

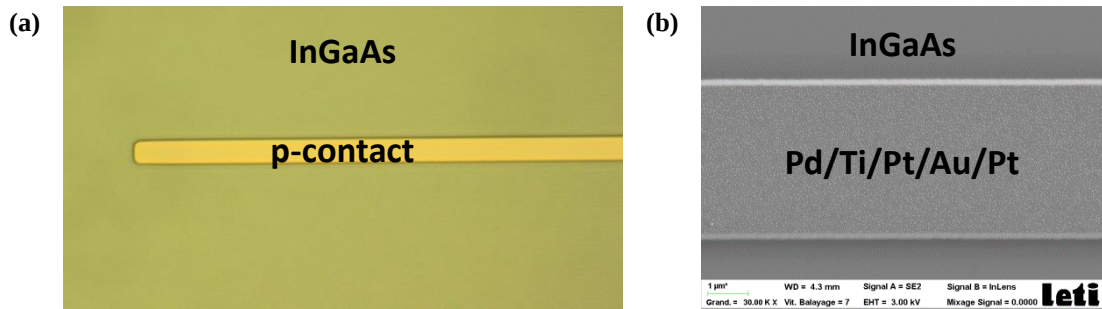


Figure III-24: Top (a) optical microscope and (b) SEM views of *p*-contact metallic layers after deposition and lift-off.

III-V waveguide patterning [see **Figure III-21(j)-(l)**]

Once the *p*-contact is deposited, the next step is to etch the III-V active waveguides. The etching process is detailed on **Figure III-25**. At first, a 400-nm-thick SiN hard mask layer is deposited, and patterned by photolithography and dry etching. The result can be seen on **Figure III-26(a)**. The III-V waveguide is then etched in three steps. At first, the InGaAs and InGaAsP contact layers are etched by using a CH₄-H₂ RIE, as shown on **Figure III-25(d)**. This etching step is not selective to InP, but it is not an issue here, since this layer will also be removed. The thick *p*-doped InP cladding is then etched by using H₃PO₄/HCl wet etching, which is fully selective to InGaAsP. Therefore, the etching is stopped by the first SCH layer. Another advantage of using wet etching is the resulting smoothness of the etched surface, which will reduce the scattering losses due to a rough surface. The main issue of this process is the fact that it is isotropic, and will also underetch the III-V waveguide. Since the etching rate depends on the crystalline orientation, the wafer orientation during bonding must be carefully chosen. In our case, the flat on the InP wafer (which can be seen on **Figure III-22 (a)**) is

parallel to the waveguides, in order to limit the etching along the width of the III-V waveguide, as shown on **Figure III-25(e)**. Finally, the SCH and MQW layers are etched down to the *n*-doped InP contact layer by using once again the CH₄-H₂ RIE, as displayed on **Figure III-25(f)**. The etching process is performed until the second SCH layer is completely removed all over the wafer. However this process is not perfectly uniform, and the edges of the wafer are etched faster than the center. Since the etching process is not totally selective to InP, when the SCH layer is fully etched at the center, the InP layer thickness will also have been etched. The thickness difference of the *n*-doped InP contact layer is estimated at 25nm between the center and the edge of the 2" stack. Therefore, the lasers at the edge of the stack are expected to have a larger series resistance than the ones at the center.

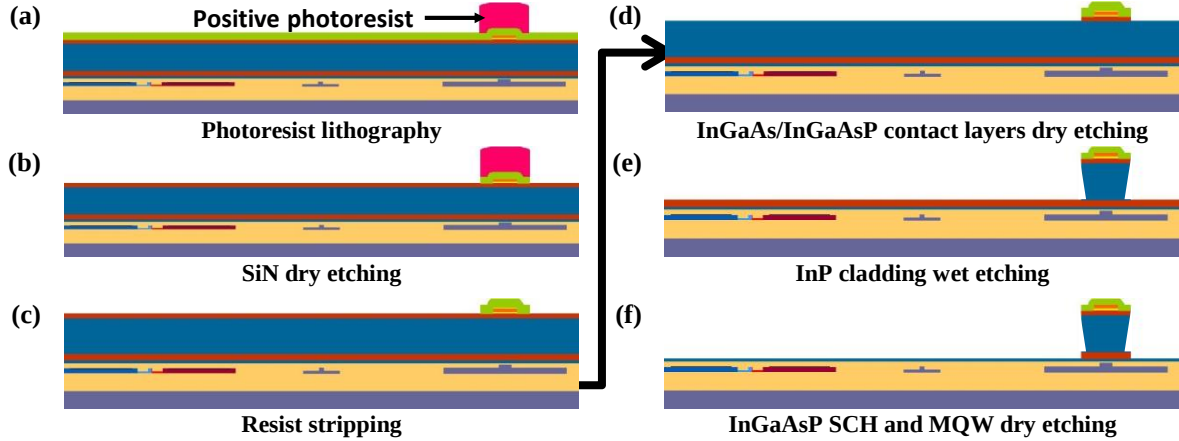


Figure III-25: III-V waveguide patterning.

Finally, the devices are separated using a resist mask and Br-based wet etching. The resulting laser structure can be seen on **Figure III-26 (b) and (c)**, with the waveguide located below the *p*-contact layer and surrounded by the *n*-doped InP layer. Now that the lasers are separated, the *p*-contact annealing is performed at 380°C during 10 seconds. No defects have been observed on the contacts after this step.

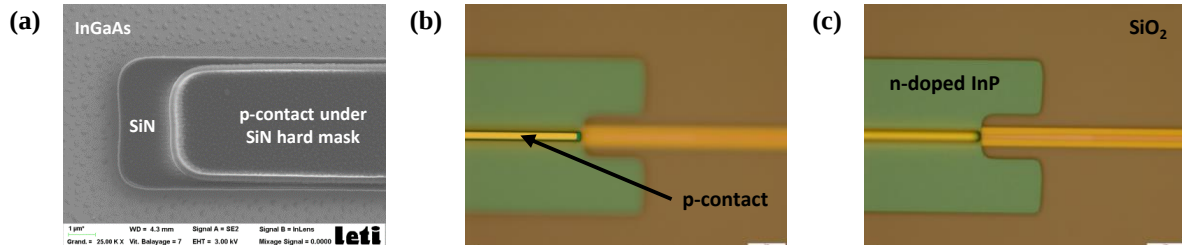


Figure III-26: (a) Patterned SiN hard mask, covering the *p*-contact metal layers. Top optical microscope views of the lasers after separation, focused on (b) the *p*-contact and (c) the *n*-doped InP layer.

Laser *n*-contact and modulator contact metallization [see **Figure III-21(m)-(n)**]

The next step is the deposition of the *n*-contact metallic layers, which are composed of Ti(30nm) /Pt(50nm) /Au(200nm) /Pt(10nm) layers. In this metallization, there is no intermixing layer, and thus an annealing step is not necessary. The same deposition and lift-off process used for the *p*-contact layer are used, leading to the structure shown on **Figure III-27**. The effects of the previously presented fabrication steps are also visible, such as the SiN hard mask covering the *p*-contact metallic layers, or the different types of etching used for the III-V waveguide. As explained previously, the width of the III-V waveguide is not much affected by the wet etching of the *p*-doped InP layer, unlike its length, which has been underetched to approximately 4-5μm. If the coupling between the III-V and silicon is sufficiently efficient, this underetch do not have a large impact, since it is located at the end of the 100-μm-long adiabatic taper, where most of the optical power has been transferred to the silicon waveguide.

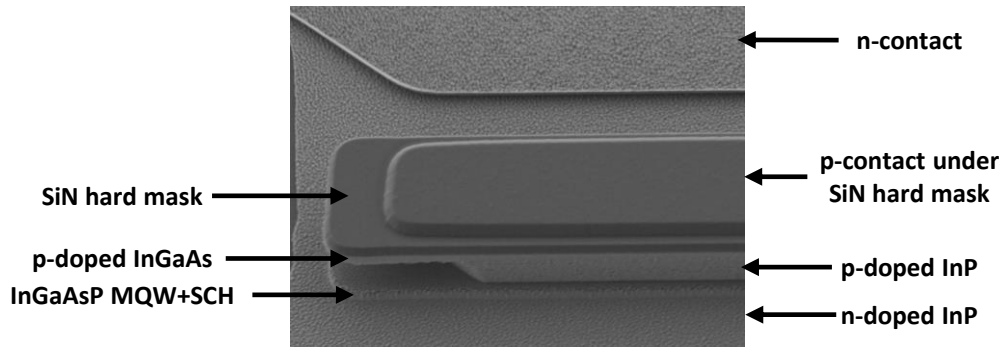


Figure III-27: SEM side view of the hybrid III-V on silicon laser after deposition of the n -contact layers.

In order to prevent the possible damages on the thin silicided region above the modulator contact region during the re-opening of the contacts (described in the next section), a thin Pt layer was deposited as a protective layer, before the thick SiN encapsulation which is the next fabrication step. The process used is illustrated on **Figure III-28**. It is similar to the one used for the metallization of the p - and n -contacts of the laser, except that the patterned photoresist used for the lift-off is also used for the opening of the silicided regions. Once the SiO_2 layer has been removed by RIE, Ni(20nm) /Ti(15nm) /Pt(10nm) layers are deposited by evaporation and patterned by lift-off. The Ni and Ti layers have been added to ensure the adhesion of the Pt layer on the silicided regions. The equipment used for the deposition being not equipped with an in-situ etching process, only ex-situ cleaning processes were used before the deposition of the metallic layers.

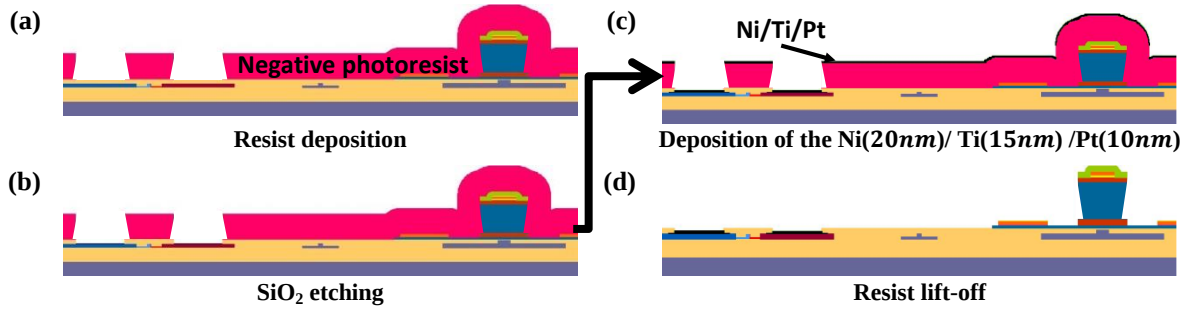


Figure III-28: Metallization of the silicon modulator.

SiN encapsulation, via opening and electrodes deposition [see Figure III-21(o)-(s)]

Once all of the metallization layers have been deposited, a 2- μm -thick SiN cladding layer is deposited to isolate the III-V waveguides. Then, the lasers contact regions are re-opened by using a photoresist mask and RIE, which is stopped by the Pt layer deposited as the last layer of the metallic contacts. The opened contacts are shown on **Figure III-29(a)**. Then, to form the laser electrodes, a 1.3 μm gold layer is deposited and patterned by lift-off, resulting in the structure displayed on **Figure III-29(b)**. The same steps are repeated for the modulator electrodes, with the same deposited thickness for the gold layer forming the electrodes.

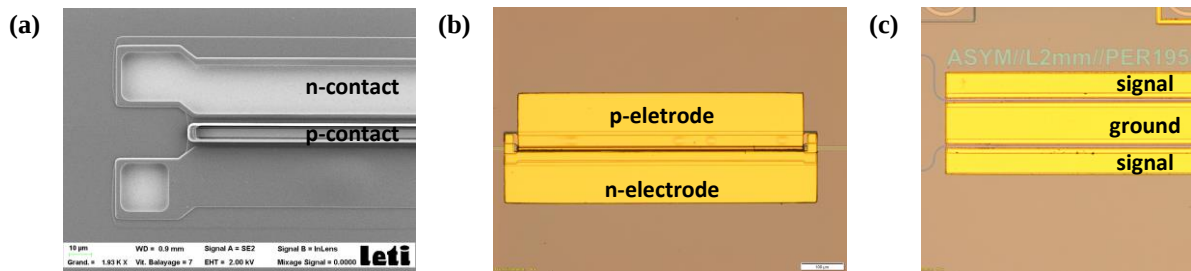


Figure III-29: Hybrid III-V on silicon laser (a) after SiN opening, and (b) after electrodes lift-off. (c) Silicon modulator after electrodes lift-off.

Heater and contact pads deposition [see Figure III-21(t)-(u)]

Finally, a 250-nm-thick NiFe layer is deposited and patterned by lift-off, in order to form the heaters. The contact pads for the heaters are then defined by the deposition of a 250- μm -thick Au layer and a last lift-off step. A top view of the completed hybrid III-V/silicon transmitter with a 2-mm-long silicon MZM is displayed on **Figure III-30**. The same figure also show several elements described earlier: a Bragg mirror, a silicon mode transformer, the patterned III-V waveguide with its metallic contact layers, the end of the transition from thick to thin silicon waveguides, a MMI, the silicon MZM waveguides, and the waveguide-to-fiber grating coupler.

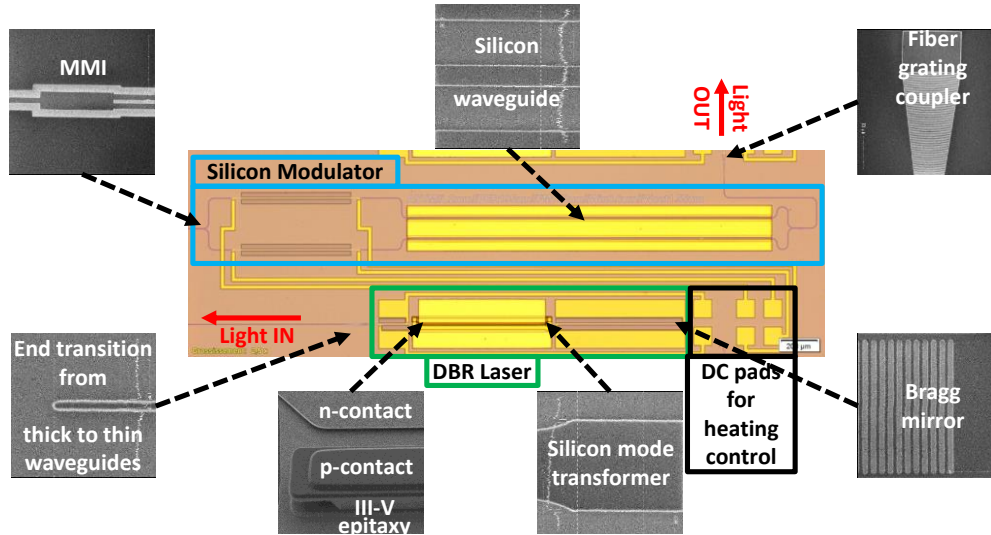


Figure III-30: Top view of the complete transmitter, with SEM views of the different components.

III-4. Integrated transmitter characterization

In this section, the characterization of the fabricated hybrid III-V on silicon high-speed transmitter is presented. First, the components forming the transmitter (passive components, DBR lasers, and silicon modulators) are optically and electrically tested via stand-alone devices fabricated along with the transmitter. The measurements realized on these components revealed several flaws in the fabrication process of the transmitter, which will help to understand the results of the complete transmitter and find the necessary improvements, though a complete de-embedding such as the one presented in **section II-4.1** is not possible. Then the characterization of the complete transmitter is presented, with static and large-signal measurements at a 25Gb/s bit rate. All the measurements presented in this section have been conducted on the same downsized III-V on SOI bonded wafer.

III-4.1. Stand-alone components

Passive components

The same testing methods and patterns as the ones presented in **section II-4.1** are used here. The die size (22mm $\times$ 22mm) is also the same. Only the components covered by the bonded III-V wafer are tested, since the etching process used for the III-V waveguide have not been optimized to be selective towards SiO₂ and silicon. Thus the non-protected components might have been damaged during the III-V etching steps, and have not been tested. Therefore, at most two components of each type were optically tested.

Unfortunately, unlike in **section II-4.1**, it was not possible to evaluate the loss of all components. Indeed, to evaluate the loss of each component, we used two testing patterns with a different number of device, or spirals with different lengths. However, not only the shortest testing pattern had more losses than in the case of the stand-alone modulator, but the transmitted power for the longest testing pattern often fell below the sensitivity of the detector, due to large optical losses. It was not possible either to extract the propagation losses in the doped waveguide through the stand-alone silicon MZMs present on the wafer either, since devices with a length

difference of 2mm have nearly the same transmission losses (between -20 and -22dB). Due to these contradictory results, it was not possible to estimate the optical losses of the components.

Due to the complex fabrication process of these transmitters and the lack of optical testing during the manufacturing, it is difficult to know which fabrication steps which might have caused these results. However, unlike in the case of the stand-alone modulator whose fabrication is presented in **section II-3**, the starting SOI layer was not 300nm , but 500nm , which was then etched to 300nm , in order to form the passive components forming the transmitter (except the ones linked to the laser as the adiabatic taper and the transition from thick to thin silicon waveguide presented in **section III-2.1**). The additional etching step might have caused a non-negligible surface roughness on top of the 300-nm -thick silicon waveguides, as well as variations on the waveguide thickness, which might explain the large losses exhibited by these components.

Thus, only two components could have been evaluated in terms of optical losses. The first is the transition from thick to thin silicon waveguide (which was not presented in **section II-4.1**, and could be used for comparison), its losses have been evaluated at approximately -0.3dB . The second is the optical grating couplers, with a maximum of transmission of approximately -5.8dB . It can be noted that since the SOI used for the transmitter has a thicker BOX than the one used in the previous chapter ($1\mu\text{m}$ instead of 800nm), a different performance was also to be expected. The transmission spectrum of the measured grating coupler is shown on **Figure III-31**.

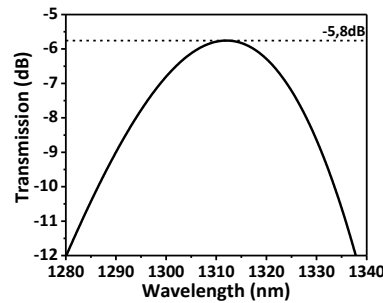


Figure III-31: Optical loss of an individual grating coupler (fitted measurement).

Hybrid III-V on silicon DBR laser

Stand-alone hybrid III-V on silicon DBR lasers are also present on the wafer for testing. Their measurement set-up is presented on **Figure III-32**. A grating coupler is connected at the end of the laser structure, which output light is collected in a single-mode fiber. The optical signal is either sent to a power-meter or an optical spectrum analyser (OSA). Since it is forward-biased, the laser is driven by the injected current, unlike the modulator which is driven by the applied voltage.

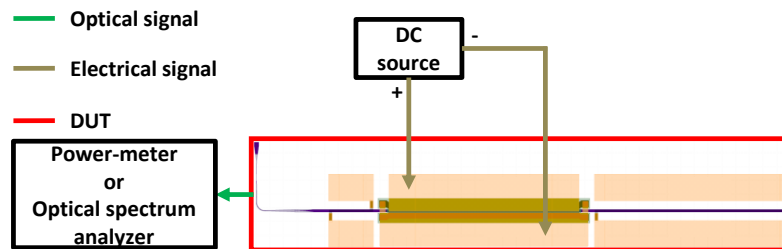


Figure III-32: Set-up for the laser electro-optical measurements.

Figure III-33 shows the performances of two DBR lasers with two different grating periods for their reflectors. DBR1 and DBR2 use Bragg reflectors with a grating period of respectively 195nm and 197nm . The $P(I_d)$ characteristic of each laser is displayed on **Figure III-33(a)**. The power in the waveguide is estimated from the grating coupler measurement shown previously. The respective current thresholds of DBR1 and DBR2 are 55mA and 49mA , equivalent to threshold current densities of 1.57kA/cm^2 and 1.4kA/cm^2 . Optical powers below 1mW are collected in the fiber, and the optical power in the waveguide is estimated to $1\text{-}3\text{mW}$ at best. These reduced output powers can be partially explained by the oxide gap between the III-V and silicon waveguides, which is lower than its optimized value (estimated at $\approx 60\text{nm}$ instead of the desired 100nm),

which in turn strongly reduces the coupling efficiency as shown on **section III-1.2**. Therefore, only a fraction of the light generated in the laser can benefit from the feedback provided by the Bragg reflectors. According to **Figure III-33(b)**, the series resistance of DBR1 and DBR2 are respectively estimated at 4.4Ω and 3.9Ω , which indicates that there was no issue with the metallization steps of the lasers.

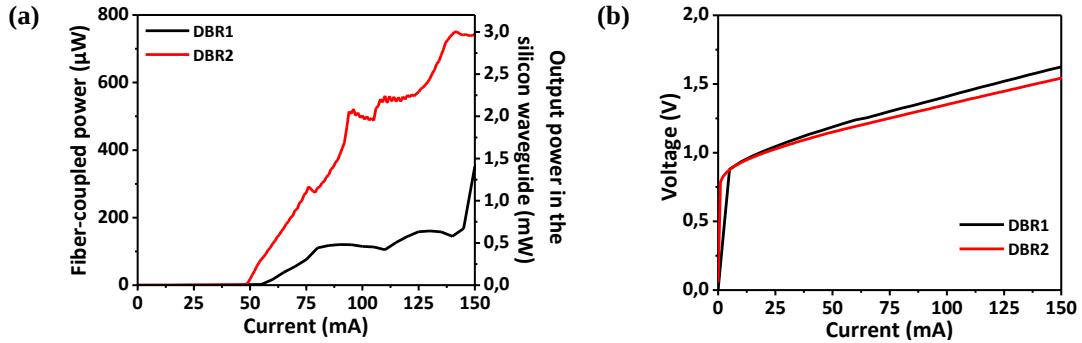


Figure III-33: (a) Optical power measured at the grating coupler output (left axis) and estimated in the silicon waveguide (right axis) and (b) voltage between the lasers electrodes, both versus pumping current.

Optical spectra corresponding to these two lasers are displayed on **Figure III-34**. Up to a certain amount of injected current (slightly above 100mA for these lasers), both lasers exhibit single-wavelength operation, with SMSRs over 30dB, as shown on **Figure III-34(a)-(b)**. Due to their different grating periods, the peak wavelengths of DBR1 and DBR2 are respectively close to 1303.4nm and 1315nm. Other modes are visible close to the lasing mode. These modes respect the phase condition for lasing, and are selected by the reflectors, but do not have enough gain to reach the threshold condition on the amplitude. The mode spacing is approximately 0.24nm, which is close to the value estimated in **section III-1.3**. When the injected current increases, some of these other optical modes have enough gain to fulfil the threshold condition, and start to lase, as displayed on **Figure III-34(c)-(d)**. Therefore, it is not possible to inject a large current in these lasers while keeping a SMSR over 30dB, which in turn limit even further their maximum output power. To solve this issue, it is necessary to reduce the number of wavelength selected by the gratings. As explained in **section III-1.3**, it would be difficult to further decrease the *FWHM* of the gratings. Thus, the solution would be to increase the *FSR*, by reducing the length of the III-V active waveguide.

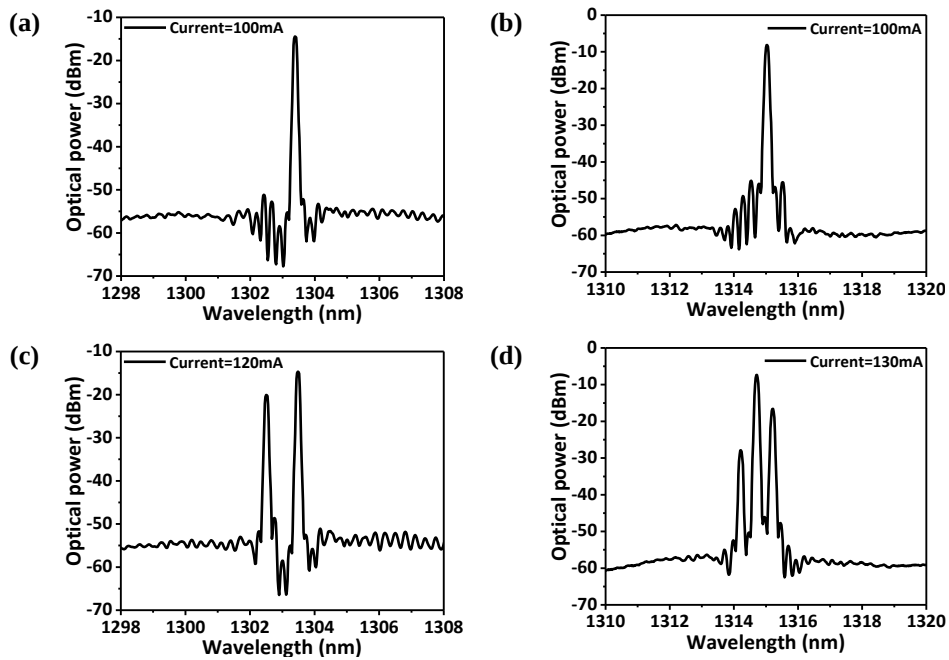


Figure III-34: Optical spectrums for (a) DBR1 at 100mA, (b) DBR2 at 100mA, (c) DBR1 at 120mA, and (d) DBR2 at 130mA.

Mach-Zehnder modulator

The static electro-optical performance of the stand-alone silicon modulators present on the wafer has been tested using the set-up described on **Figure II-32**. The FSR value was the unchanged ($\approx 4.2nm$), giving a group index $n_g \approx 4$. The average phase shift efficiency on both arms of the same MZM was evaluated and is displayed on **Figure III-35**. At a $2.5V$ reverse-bias, the phase shift is estimated to $18.9^\circ/mm$ (equivalent to $\approx 2.4V.cm$), which is close to the maximum obtained for the silicon MZMs, in **section II-4.2**.

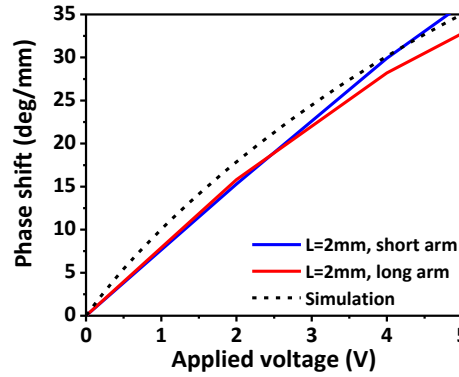


Figure III-35: Phase shift versus applied voltage, simulations and measurements for a 2-mm-long MZM.

During the measurements of the stand-alone modulators, it was noticed that if driven with a forward bias, their value of current was quite reduced compared to the stand-alone modulators presented in **chapter II**. Indeed, for the modulators co-integrated with the lasers, the measured current was at most $1mA$ for a $1V$ forward voltage, while it was above $20mA$ in the case of the modulators tested in **section II-4**. We assume this issue was caused by a native oxide layer formed between the silicided region and the contact layer deposited on the modulator, due to the lack of an in-situ etching step during deposition. This oxide layer is not thick enough to prevent the contact, but still adds a non-negligible contribution to the access resistance, and a parasitic capacitance between the electrodes and the silicon contacts of the modulators. While it was possible to reduce the resistance of this layer (and probably breaking the parasitic capacitance) by forcibly injecting a large forward current to the diode (above $400mA$), it still adds negative effects on the bandwidth of the modulator.

The small-signal electrical characterizations of the modulators are realized using the set-up presented on **Figure II-32**. The results of these measurements for each MZM length are displayed on **Figure III-36**, and compared to the results obtained in **section II-4.3**. In all cases, the measured S_{21} -parameters $-6dB$ electrical bandwidths are smaller than those obtained for the previous measurements. Below $10GHz$, the S_{11} parameters are also slightly larger for the new measurements. The trend is inverted for higher frequencies, but has a reduced effect on the bandwidth, since the signal amplitude is largely attenuated at the end of the transmission line.

As in **section II-4.3**, in order to understand the differences between the measured devices, the characteristic impedance, RF losses and RF effective index are extracted from the S-parameters of the 4-mm-long MZM. The lower value for the real part of the characteristic impedance of the new measurement below $10GHz$ causes a larger impedance mismatch with the 50Ω impedances of the VNA ports, resulting in the increase of the S_{11} parameter compared to the previous measurement. The RF effective index is also larger (up to $10GHz$), followed by a faster decrease than in the previous case. Using both **Eqs. (II.24)-(II.25)**, it can be seen that these tendencies come from the total capacitance. According to **Eq. (II.22)** the faster decrease in capacitance can be attributed to a larger access resistance, which is coherent with the assumption made of a thin oxide layer between the silicon contacts and the electrodes.

Up to $25GHz$, the losses of the modulator are also larger than for the previous case, which explains the smaller values of $-6dB$ electrical bandwidths for the S_{21} parameter. According to **Eq. (II.26)**, this trend can be either explained by larger values of line resistance or conductance. However, using **Eq. (II.23)**, it can be seen that a larger access resistance caused by a thin oxide layer would also cause the line conductance to increase. Therefore, it can be concluded that the reduced RF performances of the modulator do not come from the modifications on the electrodes shape, but on a fabrication issue during the metallization of the modulators.

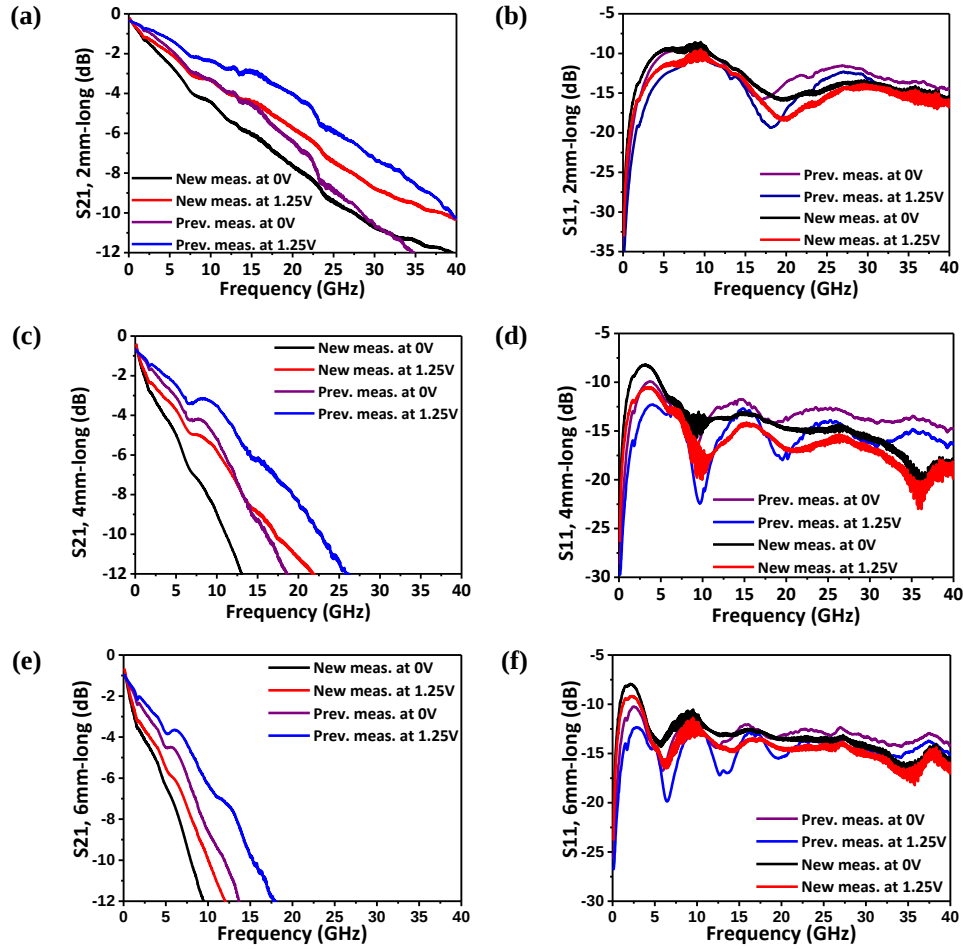


Figure III-36: Measurements of the S-parameters magnitudes for the different lengths of the active region: (a)-(b) 2-mm-long, (c)-(d) 4-mm-long, and (e)-(f) 6-mm-long. The purple and blue lines are the same as the ones presented on **Figure II-35**.

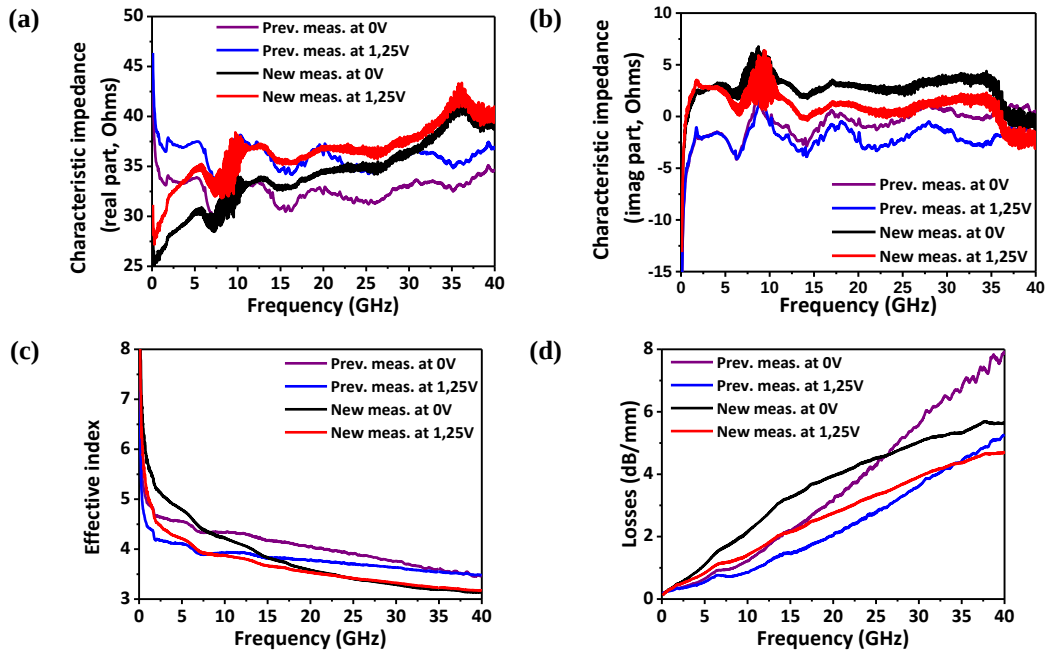


Figure III-37: Simulations and measurements of the (a) real and (b) imaginary part of the characteristic impedance, and of the (c) RF effective index and (d) RF losses.

The results of the S_{21} $-6dB$ bandwidths are summarized in **Table III-3**. Due to the larger level of optical losses of the “new” modulators, it was not possible to evaluate their modulation depth using the same set-up presented on **Figure II-37**. Nevertheless, we conducted the transmitter large-signal measurements with 2 and 4-*mm*-long silicon MZMs.

Table III-3: Measured S_{21} $-6dB$ bandwidths for the stand-alone modulators measured in **section II-4.3** and the modulators co-integrated with the hybrid III-V on silicon lasers.

Active region length (mm)	S_{21} $-6dB$ bandwidth (GHz) previously measured in section II-4.3		S_{21} $-6dB$ bandwidth (GHz) measured for the modulators	
	0V bias	1.25V reverse-bias	0V bias	1.25V reverse-bias
2	18.7	25.4	14.8	20.9
4	11	14.1	5.9	11
6	7.6	9.5	4.6	7.6

III-4.2. Complete transmitter characterization

Once the stand-alone components forming the transmitters have been evaluated, the next step is to measure the performances of the complete transmitters. At first, the static characterizations of two of these transmitters (one by peak wavelength) are shown, followed by spectral characterizations in order to study the effect of the heaters above the reflector and silicon MZM on the transmitter spectrum. Finally, eye diagrams demonstrating 25Gb/s operation for both transmitters at 0 and 10km are presented.

Static characterization

To evaluate the performances of the hybrid III-V/silicon transmitters, two of them are studied with different Bragg mirrors periods and silicon MZMs lengths. The first transmitter (referred as transmitter 1) includes a DBR laser with a grating period of 195nm and a 4-*mm*-long MZM modulator. The second one (referred as transmitter 2) uses a DBR laser with a grating period of 197nm and a 2-*mm*-long MZM modulator.

Static performances of the transmitters are displayed on **Figure III-38**. These graphs are limited to the current values for which the lasers are operating in the single-mode regime, which in turn limits their maximum output powers. During the measurements, the modulator heaters have been used to control the static phase difference between the MZM arms, and set the modulator at its maximum of transmission. Both lasers have their current thresholds at 48mA, equivalent to threshold current densities of 1.37kA/cm². The lasers series resistances are 7.1Ω for transmitter 1 and 4.5Ω for transmitter 2. The difference in series resistance comes from the non-uniformity of the dry etching step of the SCH and MQW layers, as explained in **section III-3.2**. Transmitter 1 being further from the center of the wafer than transmitter 2, its *n*-doped InP contact layer was over-etched, which explains its larger resistance value. Maximum powers collected in the fiber were respectively 37 and 52μW for transmitters 1 and 2. Considering the measured loss for the silicon MZMs from fiber-to-fiber (between 20 and 22dB), without the loss of a single grating coupler (shown in **Figure III-31**), the maximum power in the silicon waveguide can be roughly estimated between 1 and 2mW.

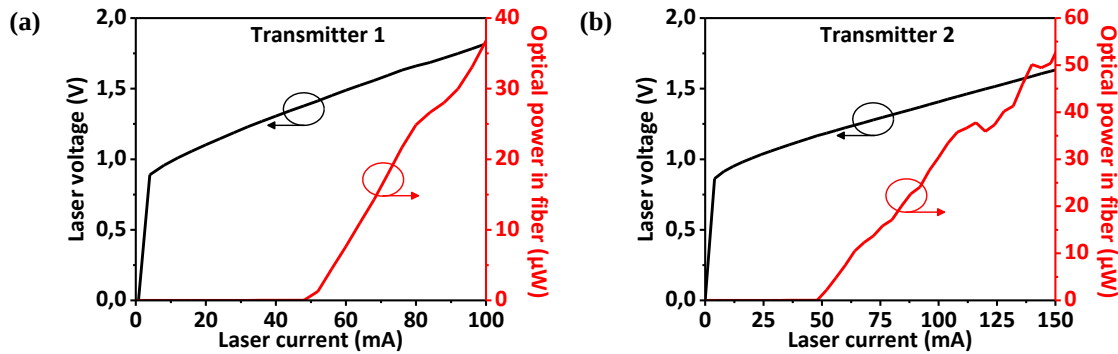


Figure III-38: Laser voltage and optical power collected in the fiber versus laser pumping current for both (a) transmitter 1 and (b) transmitter 2. For the displayed values of current, the lasers operate in single-mode regime.

Spectral characterization

Following these measurements, the effects of the heaters on the hybrid III-V on silicon transmitters are studied via an optical spectrum analyser (OSA). **Figure III-39** shows the effect of the transmitter 1 modulator heater, for a fixed laser drive of $100mA$, and at a fixed lasing wavelength. As it can be seen, the laser operates in the single-mode regime. Moreover, the interferometric spectrum of the modulator can be seen, thanks to the laser spontaneous emission over the wavelength span. By increasing the heating power on one of the silicon MZM arm and changing its phase, its spectrum can be shifted. The shift direction depends on which arm is heated. Thus, it is possible to set the laser wavelength at different working point of the modulator, such as at a maximum of transmission, at quadrature, or at a minimum of transmission. Due to the length of the undoped regions below the heaters (1mm), the electrical power needed for the heaters to create a π phase difference is approximately $50mW$. The same operation can be done with transmitter 2.

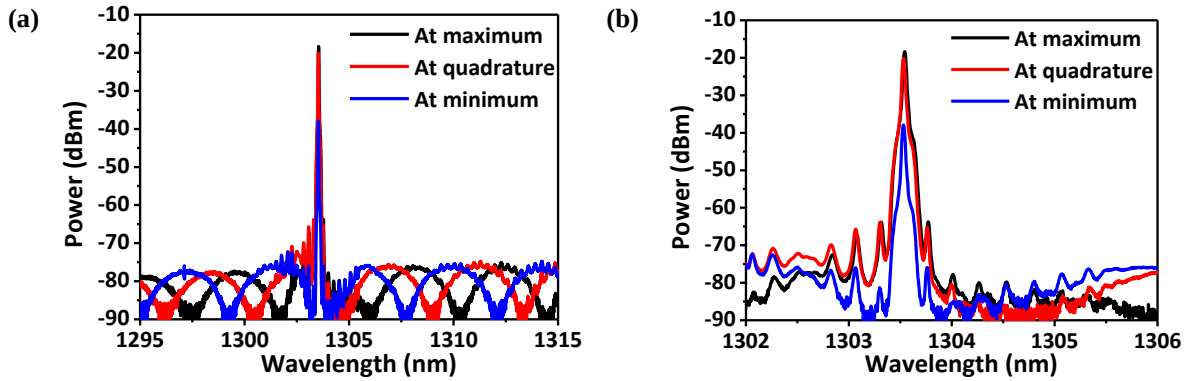


Figure III-39: (a) Broadened and (b) narrowed spectrum of transmitter 1 for a fixed laser wavelength, with different static phase difference between the MZM arms. The laser drive current is fixed at $100mA$.

Figure III-40 shows the effect of the Bragg mirrors heaters, with a fixed modulator phase for both transmitters. As explained in **section III-1.3**, by increasing the heating power on both mirrors, the mirrors spectrums will be shifted towards higher wavelengths, thus increasing the laser wavelength. The tuning is discrete, since the wavelength “jumps” between cavity modes. A large amount of heating power is required to tune the wavelength, due to the $2\text{-}\mu m$ -thick SiN layer between the heater and the silicon waveguide, leading to a reduced heating efficiency. Due to a lack of time we could not realize a spectrum with precise wavelength shifts. Thus, we focused on the maximum shift of the laser wavelength, which is up to $8.5nm$ for both transmitters. The efficiency of these heaters can be largely improved by optimizing their design, or by using a different heating technology, such as doped silicon heaters [128].

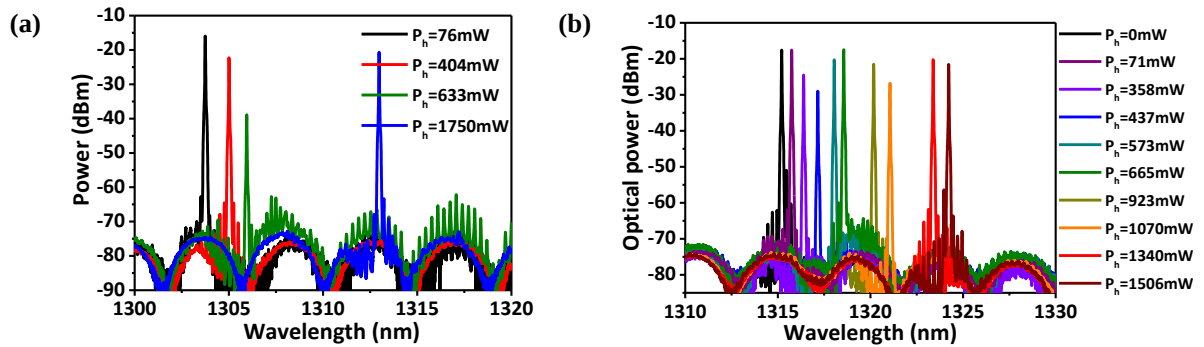


Figure III-40: Spectrum for a fixed MZM phase difference, with a tuned laser wavelength (a) for transmitter 1, with a laser drive current fixed at $100mA$, and (b) for transmitter 2, with a laser drive current fixed at $150mA$. P_h is the total electrical heating power used for both the front and back mirrors.

Large-signal characterization

Top views of a hybrid III-V transmitter under test are shown on **Figure III-41(a)-(b)**. Due to the reduced spacing between the electrical pads, and the small spacing between the probes on the testing station, it was not possible to use all the heaters at the same time. Thus, only the phase control on the modulator is used – since a single heater is enough –, while the lasers are kept at a fixed wavelength. Nevertheless, more wavelengths could have been used by operating the Bragg mirrors at the same time.

To evaluate the high speed digital performance of the transmitters, the set-up depicted in **Figure III-41(c)** is used. This set-up is nearly identical to the one shown on **Figure II-43** to test the stand-alone silicon modulators. The equipment used to drive the modulator are the same as those described in **section II-4.4**. The difference lies on the optical input and output of the device, since no external light needs to be provided. Instead, two DC-generators are used: one to drive the laser, and another one to control the modulator phase difference between its arms. Moreover, due to the output optical power limitation of the transmitters, the BOA is not suited for the measurements due to its large level of noise. Instead, an external photodetector with a better sensitivity than the optical input of our oscilloscope is used. The model is a 1544B photodetector from New Focus, which has a high internal trans-impedance gain, but which bandwidth is limited at 12GHz , making it limited for 25Gb/s measurements (which is why we did not use it for the previous large-signal measurements of the modulator). This photodetector being DC-coupled, the extinction ratio of the signal can be directly measured. As for the modulator measurements in **section II-4.4**, degradations caused by cables are de-embedded using de-emphasis, without compensating for the transmitter frequency response.

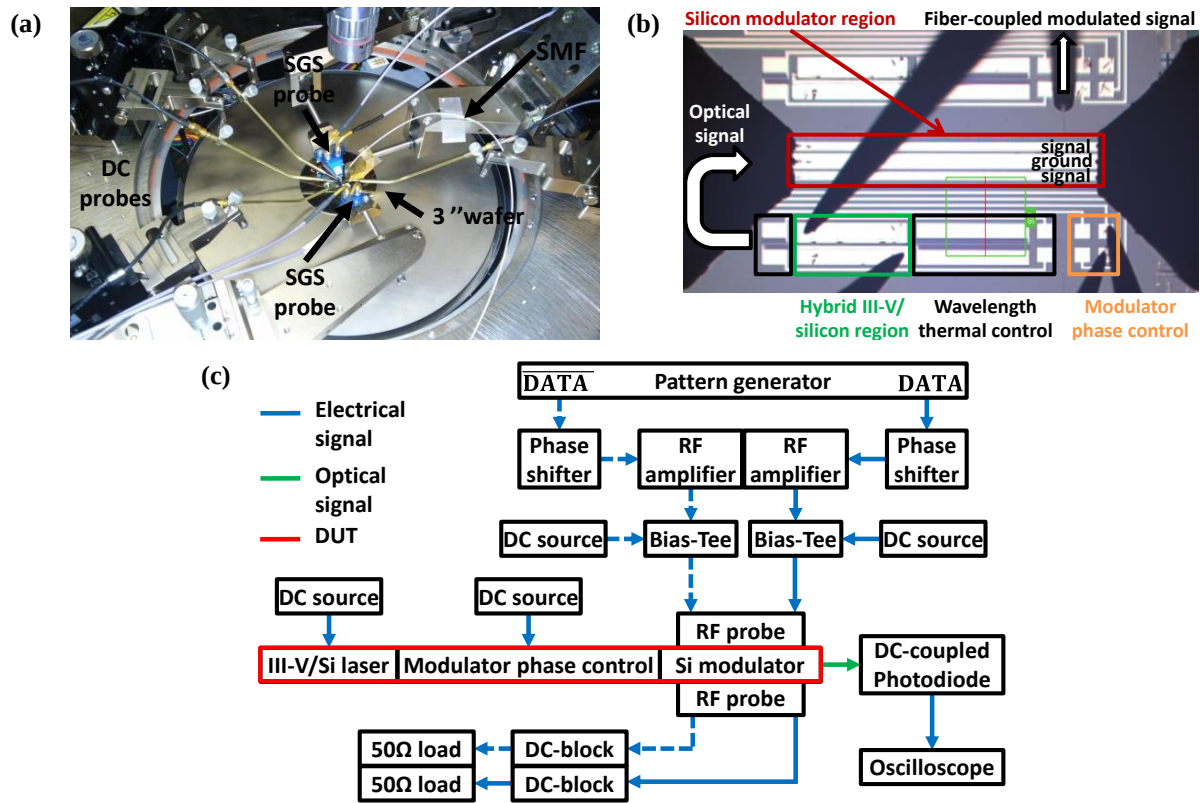


Figure III-41: (a) and (b) Pictures of the transmitter under high-speed testing. (c) Set-up for large-signal electro-optical measurements.

The resulting eye diagrams are shown on **Figure III-42** for both transmitters. As for the previous modulator large-signal measurements, a passive numerical low-pass filter (4th order Butterworth filter) with a 25GHz cut-off frequency is used to reduce the noise during the measurements. Both eyes are open at 25Gb/s , in a dual drive configuration, with $2.5V_{pp}$ send on each arm of the silicon MZM. Both arms are biased at a 1.25V reverse-bias, thus they receive NRZ signals in opposite phase, switching between 0 and -2.5V . For both measurements, the modulators are set at quadrature using the phase control heaters. For transmitter 1, the laser drive current is

set at 100mA , and its emitting wavelength is 1303.5nm . For transmitter 2, the laser drive current is set at 150mA , and its emitting wavelength is 1315.8nm . The measured extinction ratios are respectively 4.7dB and 2.9dB for the transmitter 1 and 2, which are lower than those found for the stand-alone modulator in **section II-4.4**.

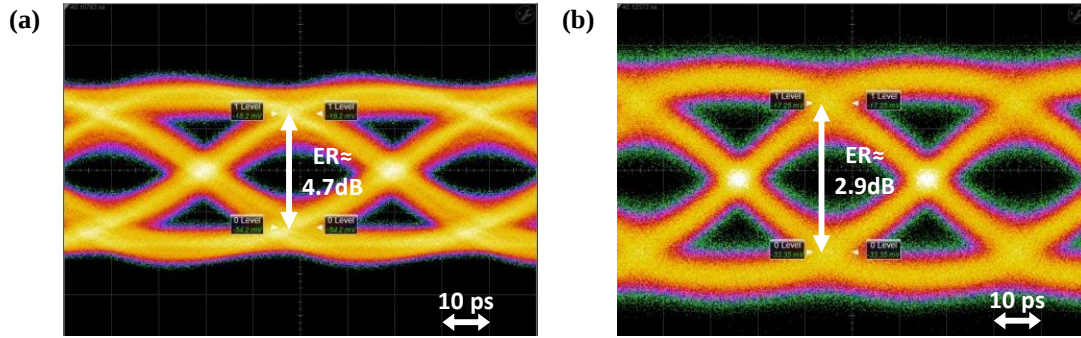


Figure III-42: 25Gb/s eye diagrams of the hybrid III-V on silicon transmitters using (a) 4-mm-long MZM at 1303.5nm (transmitter 1) and (b) 2-mm-long silicon MZM at 1315.8nm (transmitter 2), both in dual drive configuration with $2.5V_{pp}$ applied on each arm.

By comparing those eye diagrams to the ones of the stand-alone modulator on **Figure II-45(a)** and **(c)**, it can be seen that the rise and fall times are higher, which is mainly due to the photodetector bandwidth limitation. This partially explains the reduced extinction ratio of the transmitter, compared to the one of the stand-alone modulator. To verify this, we conducted the measurement on transmitter 1 with the same driving conditions, but with the BOA instead of the external photodetector as in the set-up depicted in **Figure II-43**. By maximizing the power amplification of the BOA, we were able to obtain the eye diagram depicted on **Figure III-43**. Without the bandwidth limitation of the photodetector, the extinction ratio increases from 4.7dB to 6.1dB , and the eye becomes closer to the one obtained on **Figure II-45(c)**. Nevertheless, this extinction ratio is still a bit smaller than the one found for the stand-alone modulator in **section II-4.4** (7.1dB). This is due to the bandwidth reduction of the modulators used in the transmitters and explained in **section III-4.1**, even though the static phase shift efficiencies of both modulators are equivalent.

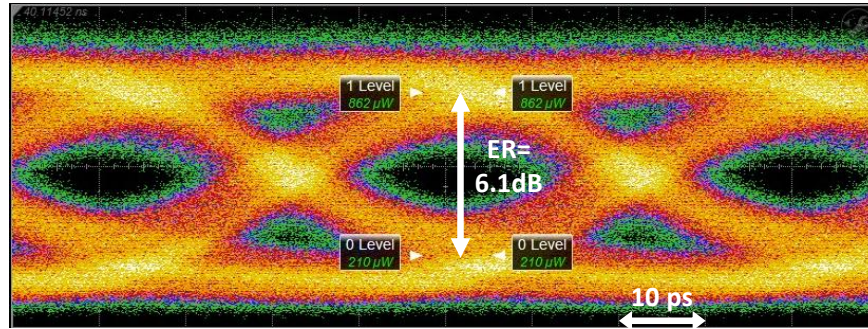


Figure III-43: 25Gb/s eye diagram of transmitter 1 in the same condition that in **Figure III-42 (a)**, but using the BOA instead of the external photodiode.

The same measurement is also realized with a 10-km-long SMF between the transmitter and the external photodetector. In this case, the fiber adds a -3.2dB optical power loss to the measurement, which renders the output power too low for the sensitivity of the external photodetector. Thus, the BOA is added with the optical filter to increase the average power over the sensitivity limit. The combination of the BOA and the optical filter provides a power amplification estimated at respectively $+5.3\text{dB}$ and $+4.4\text{dB}$ for the measurements of transmitter 1 and 2, which is enough to conduct the measurements. The resulting eye diagrams are shown on **Figure III-44**.

These eyes are nearly identical to the ones measured in the back-to-back configuration and shown on **Figure III-42**. Their extinction ratios are also the same. There are no apparent signs of dispersion, which could be expected since: 1) we used a 10-km-long SMF which prevents the modal dispersion, 2) the transmitters are

operating in the $1.3\mu\text{m}$ wavelength region, which is the region with the minimum of chromatic dispersion in the SMF (see **Figure I-1**), and 3) the silicon MZM is operated in push-pull configuration, thus its chirp is close to zero [152]. Therefore, the 10km transmission is only limited by the optical power output of the transmitter, which needs to be improved.

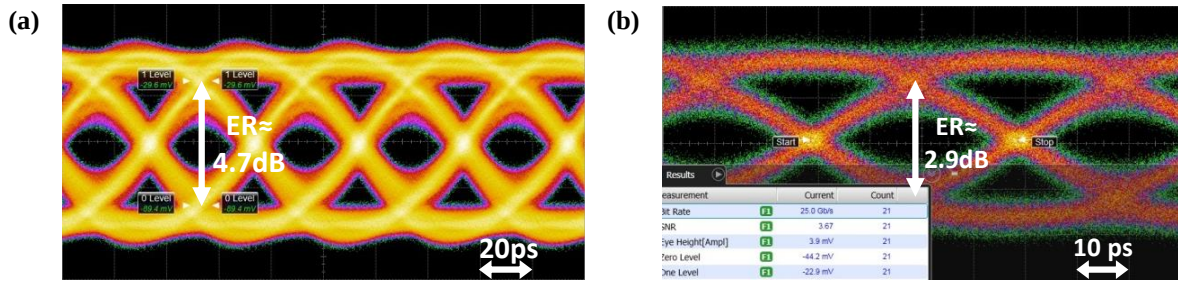


Figure III-44: 25Gb/s eye diagram of (a) transmitter 1 and (b) transmitter 2, in the same conditions that in **Figure III-42**, but with a 10-km-long SMF.

Since we used an external photodiode for the eye diagrams measurements, we did not have directly access to the values of optical power for the “ON” and “OFF” states at 25Gb/s, and we did not performed measurements to convert these values. Thus it is no possible to precisely obtain the *OMA* for these measurements, as it was the case in **section II-4.4**. However, by using the static optical power values at the maximum of transmission displayed on **Figure III-33(a)** as the “ON” state optical power and the measured extinction ratios, we can roughly estimate the *OMA*. These values are a bit overestimated, since we saw in **Table II-8** that the “ON” power level at 25Gb/s was lower than the maximum of transmission for both modulator lengths. Therefore, the *OMA* is roughly estimated at -16dBm for both transmitters in the back-to-back configuration. With the 10-km-long SMF additional losses, the *OMA* falls at -19.2dBm . Therefore, even if data transmission up to 10km at 25Gb/s is successfully demonstrated with our integrated transmitters, their performances in terms of output optical power still need to be improved.

III-5. Conclusions and improvements

We were able to demonstrate two integrated transmitters for silicon photonics, able to efficiently transmit information at 25Gb/s for distances up to 10km. These transmitters operate at two different wavelengths in the $1.3\mu\text{m}$ wavelength region, simply by changing the grating period of the Bragg reflectors used for the DBR lasers. Here, a 1nm difference on the Bragg reflectors period led to a 6.15nm difference on the laser wavelength. Additionally, since we are able to increase the laser wavelength up to 8.5nm by using the heaters above the reflectors. Thus, the transmitters can emit at any wavelength sufficiently close to the laser gain spectrum, by using a “coarse” tuning with the reflector period, and a “fine” tuning with the heaters.

As explained in **section I-4.1**, similar transmitters which co-integrate a hybrid III-V on silicon laser and a silicon modulator have only been demonstrated by two groups with bit rates below 12.5Gb/s per wavelength. In the demonstration of Alduino et al., four wavelengths fixed in the O-band, and with a 20-nm-spacing, are modulated at a 12.5Gb/s bit rate by channel [84], [85]. The four lasers output power range from 2mW to 9mW , and the modulators $V_\pi L_\pi$ product is $3.3\text{V}\cdot\text{cm}$. Thus, except for the lasers output power, the performance of our transmitter is better in terms of wavelength density, bit rate, and modulator efficiency. However, it must be noted that their demonstration was superior in terms of integration, since an Echelle grating based multiplexer was present in the circuit, and that an electrical driving circuit was also used. In the demonstration of Duan et al., a single transmitter with a tunable laser emitting in the C-band is used to modulate 8 different wavelengths at 10Gb/s [86]. The power emitted by the laser is estimated to 2.5mW at 100mA , with a SMSR of 30dB, and a tuning range of 9nm , which is similar to the performances of our lasers. The difference lies in the 3-mm-long modulator used in their transmitter, which performances are less than the ones of our modulators, with a $V_\pi L_\pi$ product of $3\text{V}\cdot\text{cm}$, and a 13GHz E/O bandwidth at a -2V bias. Moreover, this modulator could only be driven with a single arm, and required voltage swing of $7V_{pp}$, while our transmitters only need a $2.5V_{pp}$ on each arm to operate.

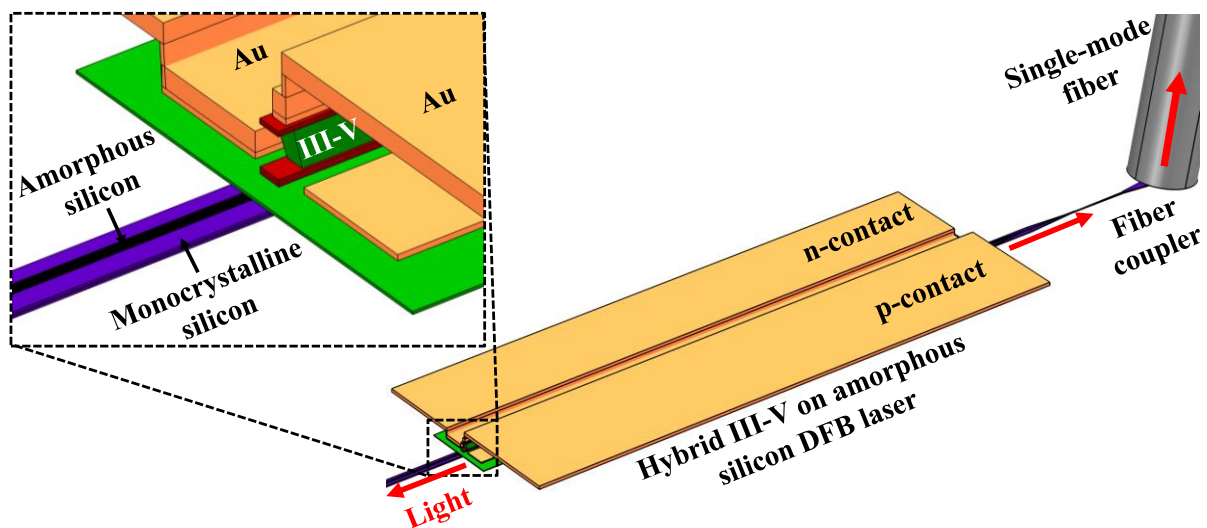
If we look at the targeted performances for our transmitter in **Table I-2**, we can see that most objectives have been fulfilled, but improvements are still needed. The targeted 25Gb/s data transmission for several wavelength has been realized, but only two wavelengths have been demonstrated, instead of four with the desired wavelength spacing of 4.5nm. Nevertheless, this objective could have been fulfilled by controlling at the same time the heaters above the Bragg reflectors and above the modulator at the same time. SMSRs above 30dB have also been demonstrated. However, the lasers could not operate in single-wavelength regime if the driving current was further increased, which limits their maximum output power. Additionally, considering **Table I-1**, larger SMSRs have already been demonstrated for this type of lasers. As explained in **section III-3.1**, this might be improved by reducing the length of the cavity. Data transmission at 10km was also realized, and even if we had to rely on an amplifying system at the output, it was seen that the eyes were not distorted, and that the optical power was the only limiting factor. Finally, even if we are quite far from the targeted OMA for now, we saw in **section II-4.4** that by reducing the losses of the passive components, it would be possible to reach the targeted OMA with output optical powers at the same level that those provided by the hybrid III-V on silicon DBR lasers demonstrated in the literature (see **Table I-1**). This assumption remains valid even by considering the additional -3.2dB optical power loss coming from the 10-km-long SMF. To summarize, the current limitations and possible improvements of our demonstrated transmitters are the following:

- 1) The optical losses increase: Compared to the performances obtained for the passive components for the stand-alone modulator demonstration in **chapter II**, their losses seem to have increased even further with the current fabrication process of the integrated transmitter. While the source of these losses is difficult to determine given the process complexity, the main hypothesis is the etching of the 500-nm-thick SOI starting layer down to 300nm to form most of the passive components. This etching step might have caused a non-negligible surface roughness on top of the 300-nm-thick silicon waveguides, as well as variations on the waveguide thickness, which might explain the increase in terms of losses.
- 2) The lasers reduced optical power: The output power of the DBR lasers was limited at 1-3mW at most, which is far from the state-of-the-art for the same type of lasers. This reduced output power is mainly attributed to the fact that the SiO₂ gap between the III-V and silicon waveguide was smaller than the one targeted in the simulations. Thus, the coupling efficiency between III-V and silicon was strongly reduced, and only a fraction of the light generated in the laser could benefit from the feedback provided by the Bragg reflectors.
- 3) The limitations of the single-wavelength operation: Even if a SMSR above 30dB was demonstrated, the lasers could not operate in single-wavelength regime for high levels of injected currents, which limited their maximum output power. This issue might be solved by reduced the cavity length of the laser.
- 4) The bandwidth reduction for the optical modulator: When co-integrated with the laser, the modulator RF performance was reduced compared to its stand-alone demonstration. This is likely due to the presence of a thin oxide layer between the silicon contacts and the electrodes during the metallization steps. This problem can be avoided by improving the cleaning process before the metallization steps.
- 5) The heaters efficiency: While they have fulfilled their role and permitted to efficiently control the modulator phase difference and the laser emitting wavelength, a large amount on electrical power was necessary to tune the laser wavelength to large values. This is due to the thick SiN cladding used during the laser process of fabrication. To solve this issue, it is necessary to improve the design of the heaters, by either getting them closer to the optical waveguide, or using a different heating technology, such as heavily doped silicon heaters

Obviously, the improvements suggested for the modulators in **section II-5** are also valid for the integrated transmitter. Amongst these issues, most of them can be corrected by improving the current fabrication process, or by improving the design of the components. However, the silicon thickness difference between the laser and the other components remains an issue which is not solved as easily as the others. Therefore, this issue is tackled in the next chapter, with a solution aiming to disturb as less as possible the current fabrication process.

Chapter IV: Hybrid III-V on amorphous silicon laser

This fourth chapter is dedicated to the demonstration of a hybrid III-V on silicon laser, based on a bi-layer made of amorphous and monocrystalline silicon. This solution is proposed to solve the SOI layer thickness incompatibility between the laser and the other components of the silicon photonic circuit. In order to verify if the solution is viable, a demonstration is realized with a distributed feed-back laser structure. The design, fabrication and characterization of this laser are detailed in this chapter. Due to the important fabrication time required for the hybrid III-V on silicon lasers, this solution has not been validated in time for the high-speed hybrid III-V on silicon transmitter, which relied on a thicker SOI waveguide. Nevertheless, the promising results obtained with this solution, and presented in this chapter can be used for future laser integrations.



IV-1. Adiabatic coupling into a thin SOI layer.....	96
IV-1.1. III-V waveguide engineering	96
IV-1.2. Silicon waveguide thickening	98
IV-2. Hybrid III-V on silicon DFB laser design.....	100
IV-2.1. Selected configuration for the DFB laser	100
IV-2.2. DFB laser structure design	101
IV-3. Hybrid III-V on amorphous Si laser fabrication	103
IV-3.1. Amorphous silicon integration schemes	103
IV-3.2. Detailed process flow.....	104
IV-4. III-V on amorphous Si laser characterization	109
IV-5. Conclusions on the use of amorphous silicon.....	111

IV-1. Adiabatic coupling into a thin SOI layer

In the previous chapter, an integrated III-V on silicon high-speed transmitter was demonstrated. However, as presented in **section III-1.1**, given the current geometry and composition of the III-V active waveguide, its effective index is too large to be efficiently coupled with the thin SOI layer standard thickness generally used for silicon photonics components. Therefore, this transmitter relies on a thick SOI starting wafer to integrate the laser, etched by RIE to reach the thickness necessary for the components constituting the rest of the transmitter (MMIs, passive waveguides, modulator waveguides and grating couplers). While this kind of fabrication process is enough for a demonstration, it is not suited for mature silicon photonics platforms [11]–[14], which are based on a thin SOI layer (equal or less than 300nm). Starting from a thicker SOI layer would induce a modification of their current fabrication processes, finely tuned to obtain high yields and performances on the standard devices. Additionally, this would add an etching process that may induce surface roughness and thickness variations which in turn will reduce the devices optical performances and reliability. This may explain why the losses of the individual components forming the transmitter (presented in **section III-4.1**) are higher than in the case of the stand-alone modulator (presented in **section II-4.1**), which has been fabricated directly on a 300-nm -thick SOI layer, without preliminary etching steps.

Rather than re-defining a new standard thickness and find a new optimal design and fabrication process for all existing silicon components, it would be more convenient to find a solution to couple the light provided by the gain region into a silicon waveguide fabricated from wafers with a thin SOI layer. In order to increase the compatibility with existing silicon photonics platforms, the additional steps for the laser fabrication should be added as late as possible in the fabrication flow of the silicon photonic circuit, so as to avoid their influence on the previously fabricated elements, because of the high temperature needed for its integration. For instance, after the formation of Ge-based photodiodes the thermal budget is limited to temperatures close to $\approx 750\text{-}800^\circ\text{C}$ [9], and after metallization steps to $\approx 400\text{-}450^\circ\text{C}$.

To solve this issue, several approaches have been proposed. The first one is to decrease the III-V waveguide effective index, by modifying its shape, and/or the composition of the epitaxial layers, to directly reach the phase-matching condition between the III-V waveguide and a thin silicon waveguide [60], [153], [154]. The second one is to locally increase the effective index of the silicon waveguide by using epitaxial growth, in order to form a coupling section with a thicker silicon region and transfer the light into a thin waveguide [144]. While these solutions gave promising results, they also present flaws in the laser performances, fabrication or with their integration in the process flow of the silicon photonic circuit due to large thermal budgets. Finally, we propose a solution based on the deposition of an amorphous silicon layer, which is compatible with the fabrication of the photonic circuits, and enable to form efficient hybrid III-V on silicon laser structures.

IV-1.1. III-V waveguide engineering

Up to now, the solution considered to couple light from the III-V gain region into the silicon waveguide has been to decrease/increase the width of the silicon waveguide, for a given isolated III-V waveguide refractive index, in order to invert the ratio of effective indices between the isolated silicon and III-V waveguide along the propagation direction. Thus, if the isolated III-V waveguide effective index would also be reduced, the ratio inversion should be reached for a thinner silicon waveguide, leading to an efficient coupling.

A simple way to reduce the effective index of the isolated III-V waveguide is to reduce its width. By using an inverted double-taper architecture in both III-V and silicon waveguides (see **Figure IV-1(a)**), a hybrid III-V on silicon DFB laser has been realized with a 400-nm -thick silicon waveguide. The III-V taper was ended with a 500-nm -wide tip, fabricated by 300nm UV contact lithography and wet under-etching of the III-V waveguide down to the n -contact layer, as shown on **Figure IV-1(b)**. Using this taper, the laser showed excellent performances, displayed on **Table I-1** (ref. [60]). Therefore, by decreasing further the width of the III-V waveguide, it should be possible to couple light in even thinner silicon waveguides. For silicon waveguides below 300nm , it would require III-V tips narrower than 300nm [70], [155], [156]. However, this solution is limited by the process of the III-V stack. Indeed, given the thickness of the III-V stacks (generally between 2 and $3\mu\text{m}$), III-V tips narrower than 500nm are quite challenging to fabricate in a repeatable way using standard UV contact lithography [59], [70].

Recently, a hybrid III-V on silicon FP laser has been realized with a 300-nm-thick silicon waveguide [153]. It relies on double-taper architecture in both III-V and silicon waveguides in the same direction (see **Figure IV-1(c)**). The III-V taper was ended with a 150-nm-wide tip, fabricated by electron beam lithography, and etched down to the bonding layer, as shown on **Figure IV-1(d)**. However, it should be noticed that the III-V epitaxy did not have a SCH layer between the MQW region and the n -contact layer, which also reduced the confinement in the quantum wells, and simplified the optical coupling in the silicon waveguide. Additionally, the laser operated only up to 20°C in CW regime.

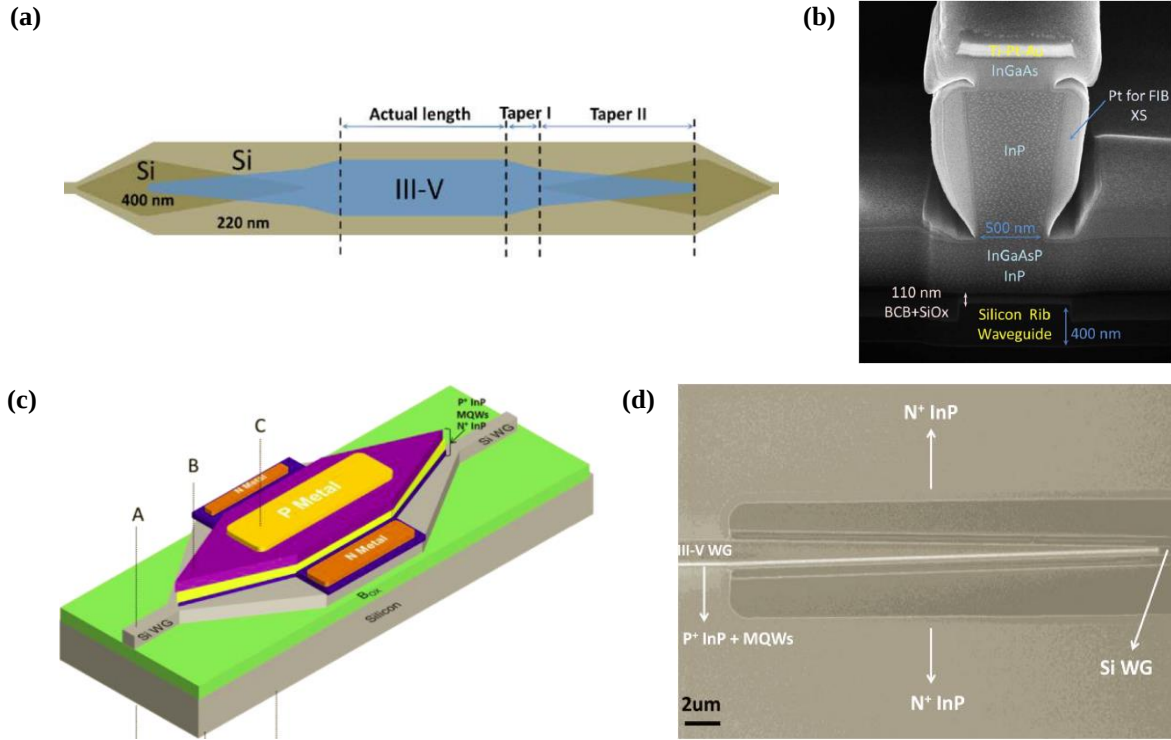


Figure IV-1: (a) Schematic view of an inverted double-taper architecture, and (b) scanning electron microscope image of its III-V tip [60]. (c) Schematic view of a same direction double-taper architecture and (d) scanning electron microscope image of its III-V tip [153].

The shape of the III-V waveguide can also be modified in a different way to reduce its effective index. By etching a wedge in the top layers and undercut the MQW region down to 800 nm, the effective index of the III-V waveguide is effectively reduced and can be coupled to thin SOI layers. An electrically pumped laser in pulsed regime was obtained with this structure [154], showed on **Figure IV-2**. However this type of waveguide is quite complex to process and the structure was also reported to be mechanically instable. This laser was operated only in electrically-pulsed regime.

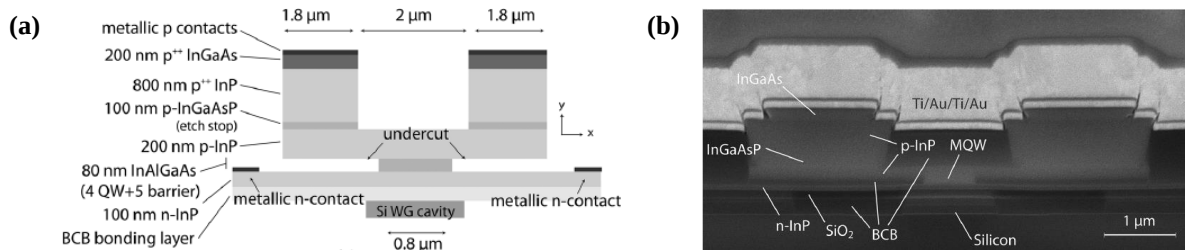


Figure IV-2: (a) Schematic view and (b) scanning electron microscope image of the III-V waveguide cross-section used in [154].

To summarize, while it might be possible to couple light from the gain region into a thin silicon waveguide by reducing the effective index of the III-V waveguide, this method relies on difficult processing of the III-V stack. The width of the waveguide must be narrowed down to values which are difficult to reach by standard UV

lithography, due to the topography of the III-V region. This issue might be solved by improving the processing of the III-V layer, but it also goes against the principles of silicon photonics which are to transfer the maximum of functions and fabrication processes in the silicon, to benefit from its excellent fabrication maturity. Modifying the III-V stack composition (such as removing SCH layers) to improve the coupling is also a possibility, but results in a trade-off on the laser performances.

Therefore, we decided to keep the current approach, without modifying the large III-V waveguide, and to keep all functions – except for optical gain – and fine processing steps in the silicon region. This way, no trade-offs are required for the laser, and the processing steps of the III-V region do not increase in complexity.

IV-1.2. Silicon waveguide thickening

If the III-V waveguide is not modified, the remaining solution to reach the phase-matching condition with a thin SOI layer is to locally increase the thickness of the silicon waveguide below the III-V region, in order to increase its effective index. Two solutions exist to increase the thickness of the SOI layer: epitaxial growth and deposition. A hybrid III-V on silicon laser based on epitaxial growth has been demonstrated in [144]. Its fabrication starts from a SOI wafer with a 300-nm-thick silicon layer. After the patterning of the thin waveguides, the thickness of the waveguides used for coupling with the III-V region is increased up to 600nm by epitaxial growth of silicon. After planarization, a III-V stack is bonded to form the laser structure. While this solution provides thick silicon waveguides without modifying the silicon thickness of other existing components, it may be difficult to implement in the later steps of a complete flow, due to the high thermal budget required for the silicon epitaxial growth.

The other solution, which has not been previously demonstrated in the literature, is to deposit a layer of polycrystalline or amorphous silicon. Amongst the two of them, polycrystalline silicon presents the advantage to be largely superior as an electrically active material [105]. Therefore, it is already used to form carrier-accumulation silicon modulators [112], [114] as explained in **section II-1.4**. Its main issue is the important propagation losses ($> 70\text{dB/cm}$ equivalent to a material absorption loss $\alpha > 45\text{cm}^{-1}$) caused by small grain size – resulting in lots of grain boundaries –, as well as the surface roughness of the waveguides formed with as-deposited polycrystalline silicon [157], which would be detrimental for a laser. To reduce the losses, the polycrystalline silicon can be deposited at low temperature ($\approx 550^\circ\text{C}$) to obtain a smooth surface, and annealed at high temperature ($\approx 1000\text{-}1050^\circ\text{C}$) to increase the size of the grains, resulting in propagation losses on the order of $\approx 10\text{dB/cm}$, or material absorption coefficient of $\alpha \approx 7\text{cm}^{-1}$ [158]. However, the losses remain large, and the high temperature needed to reduce the losses makes it difficult to integrate in the later steps of fabrication.

On the other hand, amorphous silicon can be deposited at temperatures below 400°C , and does not need a high temperature annealing to exhibit low propagation losses. A schematic view showing the differences between crystalline and amorphous silicon is shown on **Figure IV-3**. In the case of monocrystalline silicon, each atom is covalently bonded to four neighbouring atoms, and a unit cell can be defined to describe the full crystal lattice. In the case of amorphous silicon, while most silicon atoms still have covalent bonds with four neighbours, the ordered structure has been lost, and becomes a continuous random network. This results in weaker bonds, which can be broken and lead to coordination defects, where silicon atoms are bonded to three atoms and have one unpaired electron, referred to as dangling bond. In pure amorphous silicon, the concentration of coordination defect is up to $10^{21}\text{defect/cm}^3$, so 2% of the total number of atoms in silicon ($5 \times 10^{22}\text{atoms/cm}^3$). The dangling bonds are the main source of optical losses in amorphous silicon. To reduce their density, they can be passivated by hydrogen during the deposition, as displayed on **Figure IV-3(b)**. This way, hydrogenated amorphous silicon (or a-Si:H) is formed and the defect density is reduced to $10^{15}\text{-}10^{16}\text{cm}^{-3}$ [159]. Due to its reduced thermal budget, a-Si:H has been considered for a long time to form passive waveguides after the metallization steps, for multi-level photonic circuits [160]. Material propagation losses (or absorption coefficient) as low as 1.6dB/cm (0.36cm^{-1}) and 0.5dB/cm (0.11cm^{-1}) have been respectively reached in the $1.3\mu\text{m}$ and $1.55\mu\text{m}$ wavelength region [161]. a-Si:H nanowire waveguides (with a $\approx 500 \times 220\text{nm}^2$ cross-section) with propagation losses below 3.5dB/cm in the $1.55\mu\text{m}$ wavelength region have been demonstrated several times [162]–[165], which is comparable to their monocrystalline counterparts with similar dimensions. Additionally, as in the case of monocrystalline waveguides, the losses of these waveguides are dominated by the sidewall roughness, and not the a-Si:H absorption coefficient [163], [164].

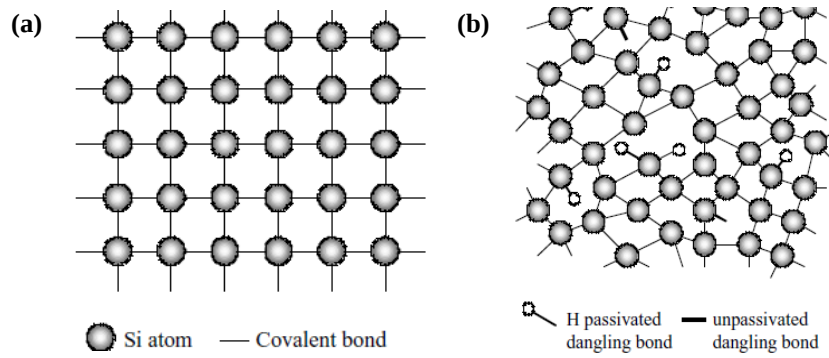


Figure IV-3: Schematic view of the atomic structure of (a) monocrystalline silicon and (b) hydrogenated amorphous silicon. Images from [159].

While as-deposited a-Si:H waveguides present low propagation losses, great care must be taken with the thermal budget following the deposition. Even if the amorphous silicon layer is deposited at 400°C , the optical losses can start to increase slowly if the waveguides are annealed at temperatures above 350°C , and exponentially above 400°C [162], [163]. These losses originate from the breaking of the Si-H bonds at increased temperature, and out-diffusion of the hydrogen atoms, leaving the dangling bonds free to absorb light again. While the annealing time does not seem to have a large influence on the losses [162], it is not the case for the annealing atmosphere. If it contains hydrogen atoms to passivate the newly formed dangling bonds, it has been demonstrated that the optical losses will be lower than an annealing without hydrogen atoms at the same temperature [163]. Therefore, it has been shown that a SiO_2 encapsulation at 400°C realized after the a-Si:H deposition will not increase its losses, due to the hydrogen atoms provided by the SiH_4 dissociation during the SiO_2 layer deposition [163], [164].

Given its advantages in terms of deposition temperature – making it compatible with late integration in the process flow – and its optical losses comparable to monocrystalline silicon, the deposition of an a-Si:H layer is chosen to locally increase the effective index of the silicon waveguide. Therefore, a waveguide with the same geometry as the one used for the transmitter, but with a 300-nm-thick SOI layer as the slab section and a 200-nm-thick a-Si:H layer as its rib section – as shown on **Figure IV-4(a)** – should enable an efficient light coupling with the III-V waveguide. We do not precisely know the material absorption coefficient of the a-Si:H layer used during the fabrication of the laser. Nevertheless, the same deposition conditions have been used in another work to form a-Si:H nanowire waveguides (with a $\approx 500 \times 220\text{nm}^2$ cross-section), and has shown propagation losses of 4.5dB/cm , close to the state-of-the-art [166].

Figure IV-4(b) shows the simulated confinement factor of the even supermode, in the rib and slab sections of the silicon waveguide based on a-Si:H and monocrystalline. It can be seen that at phase-matching, the confinement in the a-Si:H region is close to 16%, and stays below 33% when the silicon waveguide width increases. Therefore, the influence of the material absorption losses will be even more reduced than in the case of the a-Si:H nanowire waveguides.

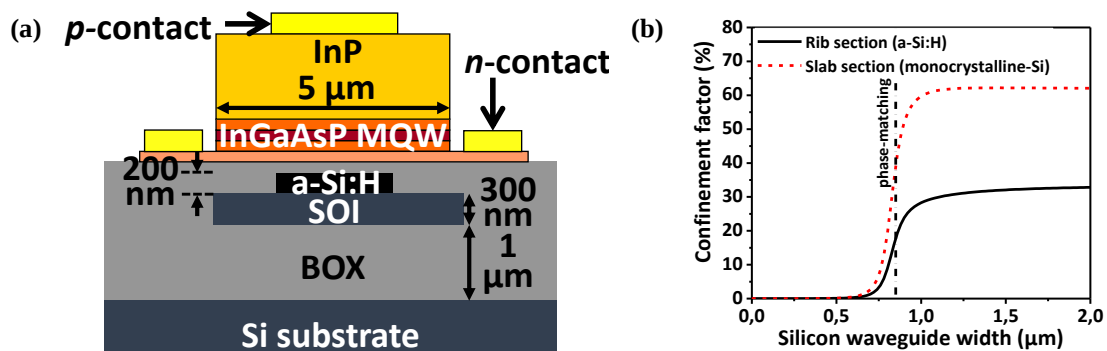


Figure IV-4: (a) Schematic cross-section of the laser based on a-Si:H and monocrystalline silicon. (b) Simulated confinement factor of the even supermode in the rib and slab sections of the silicon waveguide.

IV-2. Hybrid III-V on silicon DFB laser design

In order to demonstrate a hybrid III-V on silicon laser based on a bi-layer made of amorphous and monocrystalline silicon, a distributed feed-back structure is chosen, since there is no need for wavelength tuning in this case. In this section, the general principles of operation of the DFB laser are presented, followed by the design of the laser structure used for the demonstration.

IV-2.1. Selected configuration for the DFB laser

Unlike the DBR laser presented in the previous chapter, where two gratings are located at each side of the active gain region, in the case of a DFB laser, a grating is distributed along the active gain region. The conditions to reach laser operation can still be applied in the same form to these structures, with the gain in the complex mirror reflection coefficients (r_{g1} and r_{g2}), and without passive regions.

A schematic view of a standard DFB structure with a single grating is presented on **Figure IV-5(a)**. As it was seen in **section III-1.3**, the period of the grating is connected to the Bragg wavelength by the relation $\Lambda = \lambda_g / 2\bar{n}$, where $\bar{n}$ is the average value of the effective indices of the non-etched and etched cross-sections of the grating. In the case of a standard DFB laser with a uniform grating, the cavity length can be taken anywhere within the grating, and is equal to half the grating period: $L_a = \Lambda / 2 = \lambda_g / 4\bar{n}$. Thus it is said to be a quarter-wavelength long. In this case, there is no solution for the lasing condition at the Bragg wavelength, and the DFB is said to be antiresonant. The entire gain spectrum must be scanned in order to find which modes satisfy the lasing amplitude and phase conditions. If there are no perturbing reflections, two modes (which are equally spaced on each side of the Bragg wavelength) will reach the threshold condition simultaneously, leading to the spectrum displayed on **Figure IV-5(a)** [3]. These two modes actually correspond to the edges of the stop-band on the grating reflectance spectrum, and are referred as the first side-modes. Other side-modes located further from the Bragg wavelength can also fulfil the phase threshold condition. However, since they benefit from a lower reflection (or a larger mirror loss) than the first two side-modes, they need more gain to reach the amplitude threshold condition. Therefore, they are suppressed by the first two side-modes and become non-lasing side-modes. The gain difference necessary to reach the amplitude threshold condition between the modes is known as the gain margin.

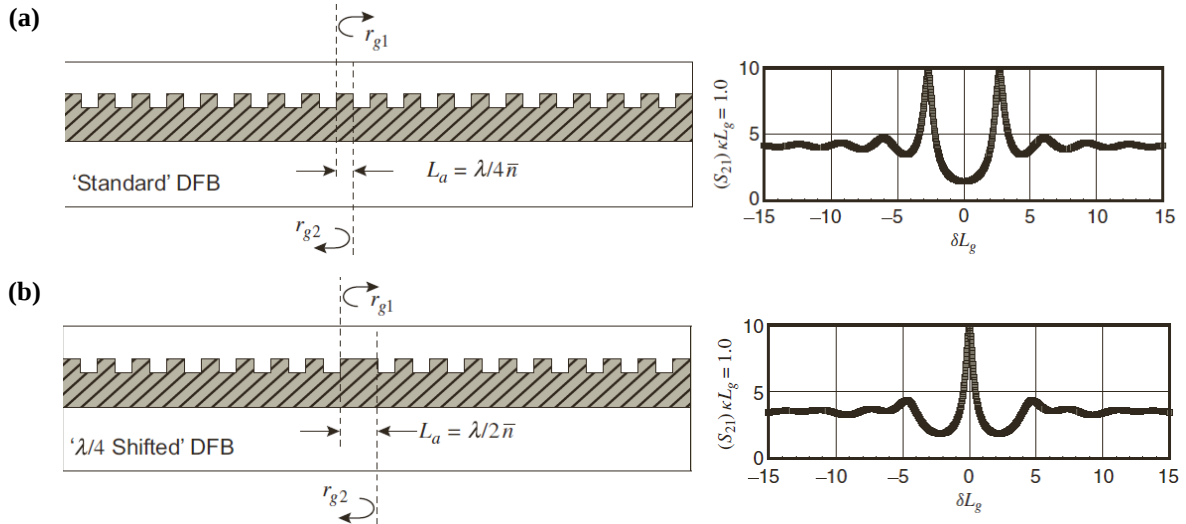


Figure IV-5: Schematic transversal view and spectra of (a) a standard DFB laser and (b) a quarter-wavelength shifted DFB laser. The entire grating is distributed along the active gain material. δ is the same detuning parameter from the Bragg wavelength (λ) defined in **Eq. (III.7)**. Images from [3].

While it is possible to suppress one of the unwanted modes by adding reflections at one side of the DFB, which will generate a gain margin between the first two side-modes, this method does not guarantee single-mode operation. An easier way is to create a “defect” in the grating, by introducing an additional section in the grating. Generally, a section with a length $\lambda_g / 4\bar{n}$ is added at the center of the grating, and devices based on this

additional section are known as quarter-wavelength shifted (QWS) DFB lasers. A schematic view of a QWS DFB structure can be seen **Figure IV-5(b)**. In this case, the cavity length is equal the grating period: $L_a = \Lambda = \lambda_g/2n$, and is said to be a half-wavelength long. Unlike the case of the standard DFB, the laser is resonant at the Bragg wavelength. Even in this case, side-modes located at each side of the Bragg wavelength can still fulfil the phase threshold condition. Nevertheless, since the mode at the Bragg wavelength has the highest reflectance (or the lowest mirror losses), it is the first to reach the threshold condition, and effectively suppress the side-modes, thus leading to the single-wavelength spectrum displayed on **Figure IV-5(b)**.

Several designs exist for the phase-shifted section which does not necessarily need to be a quarter-wavelength or even to be located at the center of the DFB to reach the single-mode regime. However, DFB lasers with a quarter-wavelength shifted section located at the center of the grating remain the most commonly used structures, which guarantee that the lasing mode will be at the Bragg wavelength and offer the maximum SMSR [3], [167]. Therefore, in order to realize our demonstration of hybrid III-V on amorphous silicon, we focused our attention on these types of DFB lasers.

IV-2.2. DFB laser structure design

A schematic view of the hybrid III-V on amorphous silicon QWS DFB laser is presented on **Figure IV-6**. Its cross-section is the same as the one displayed on **Figure IV-4(a)**. The III-V active waveguide is similar to the one presented in **section III-1.1**, and based on InGaAsP MQW which maximum gain is centred around $1.3\mu\text{m}$. The MQW layers are surrounded by two InGaAsP SCH layers, and two *n*- and *p*-doped InP layers. The III-V stack and the a-Si:H layer are separated by a 100-nm-thick oxide gap. The III-V active waveguide is $700\text{-}\mu\text{m}$ -long and $5\text{-}\mu\text{m}$ -wide. Therefore, by taking into account the $100\text{-}\mu\text{m}$ -long adiabatic tapers in the silicon waveguide and located at each side of the III-V waveguide, the total grating length is fixed at $500\mu\text{m}$. The grating is realized within the a-Si:H layer, with a quarter-wavelength shifted section located at its center.

One of the laser output is connected to the transition from thick-to-thin silicon waveguide presented in **section III-2.1**. Thus, other silicon photonic components can be defined in the thin monocrystalline silicon waveguide. In our case, a surface-to-fiber grating coupler is used to couple the light in a SMF for measurements. The other output is connected to a hybrid III-V on silicon photodiode based on the same III-V stack than the laser (not represented on **Figure IV-6**). Unfortunately, these lasers were ready for characterization at the end of this work, and we did not have the time to evaluate the output signal in the silicon waveguide with this photodiode. Therefore, since the QWS DFB structure is symmetric, we make the assumption that the same optical power come out of both laser outputs.

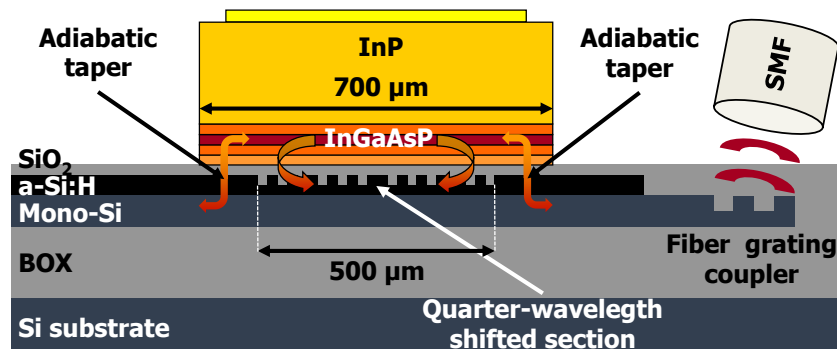


Figure IV-6: Schematic longitudinal view of the laser based on a-Si:H and monocrystalline silicon layers.

The next step is to fix the geometrical parameters for the grating. Here, the total length of the grating and the duty cycle are respectively fixed at $500\mu\text{m}$ and 0.5. Thus, the parameters left are the grating period, the etching depth, and the width of the rib section below the active III-V waveguide. Unlike the case of the DBR where the width was fixed at $3\mu\text{m}$, the grating is now below the gain region, and such a large value would result in a weak confinement of the even supermode in the MQW, as shown previously on **Figure III-4(b)**. Therefore, in order to keep a modal confinement over 10% in the MQW layers, the grating width must remain below $0.8\mu\text{m}$. In order to fix the geometrical parameters of the grating, the same equations defined in **section III-1.3** are used. The main

difference is that since the grating is now part of the hybrid III-V on silicon waveguide, the effective index to consider in the calculations is the one of the even supermode.

We have seen previously that for a QWS DFB, the cavity length is a half-wavelength long. Thus, for a targeted Bragg wavelength of $1.31\mu\text{m}$, it will be negligible compared to the grating length. Since the quarter-wavelength shifted section is located at the center of the grating, the peak reflectance on each side of the cavity can be evaluated using **Eq. (III.10)** with a length equal to $250\mu\text{m}$ (half the complete grating length). The result is shown on **Figure IV-7(a)**. It can be seen that in order to reach a peak reflectance between 10 and 90% on each side of the cavity, the grating coupling constant should be between 6 and 72cm^{-1} .

As shown by **Eq. (III.6)**, the grating coupling constant is proportional to the difference $\Delta\tilde{n}$ between the effective index of the non-etched and etched cross-sections of the grating. Here, $\Delta\tilde{n}$ depends on both the silicon rib waveguide width below the III-V region and the grating etching depth. The effective indices of the even supermode have been calculated by FEM for several values of widths and etching depths. A map of the grating coupling constant versus these two parameters is displayed on **Figure IV-7(b)**. As it can be seen, the grating coupling constant mainly depends on the waveguide width rather than on the etching depth, which effect strongly decreases above 60nm . Indeed, the grating coupling constant depends on the mode overlap on the etched region. As shown previously on **Figure III-4(d)**, when the silicon waveguide width is below $0.6\mu\text{m}$, the even supermode is weakly confined in the silicon waveguide. Thus, only a small fraction of the mode overlaps with the grating, and the grating constant is small. When the silicon waveguide width increases, the even supermode confinement in the silicon waveguide increases drastically, and so does the grating coupling constant. For a fixed silicon waveguide width, the grating coupling constant will first increase with the etching depth. However, since the mode overlap on the etched part of the grating is fixed by the silicon waveguide width, for widths below $0.8\mu\text{m}$, the even supermode is still mainly confined in the III-V waveguide, and the value of grating coupling constant will eventually stop to increase.

According to the result of the calculations displayed on **Figure IV-7(b)**, for etching depths above 60nm , silicon waveguide widths between 0.6 and $0.8\mu\text{m}$ lead to grating coupling constants between 6 and 40cm^{-1} . These relatively low values are due to the 100-nm -thick SiO_2 gap between the grating and the III-V waveguide, which strongly reduces the fraction of the even supermode overlapping with the grating. Finally, the width of the grating is fixed at $0.8\mu\text{m}$, and the etching depth at 75nm , leading to a coupling constant of 40cm^{-1} . The last parameter to fix is the period of the grating, in order to target a Bragg wavelength close to $1.31\mu\text{m}$. Given the geometry chosen for the grating, its average effective index is approximately ≈ 3.3 , which leads to a grating period close to 198nm .

It can also be noted that unlike the case of the DBR laser discussed in the previous chapter, the odd supermode can also profit from the feedback provided by the grating, and reach the threshold condition. Indeed, as it can be seen on **Figure III-4(d)**, for a waveguide width of $0.8\mu\text{m}$, the odd supermode is even more confined in the silicon waveguide than the even supermode, while keeping a 6% confinement in the MQW. Thus for the same grating geometry, the coupling constant is estimated at 197cm^{-1} , leading to a peak reflectance at each side of the cavity close to 100% (according to **Figure IV-7(a)**). Nevertheless, due to the adiabatic tapers, the mode would stay confined in the III-V waveguide, thus limiting the optical power coupled in the fiber.

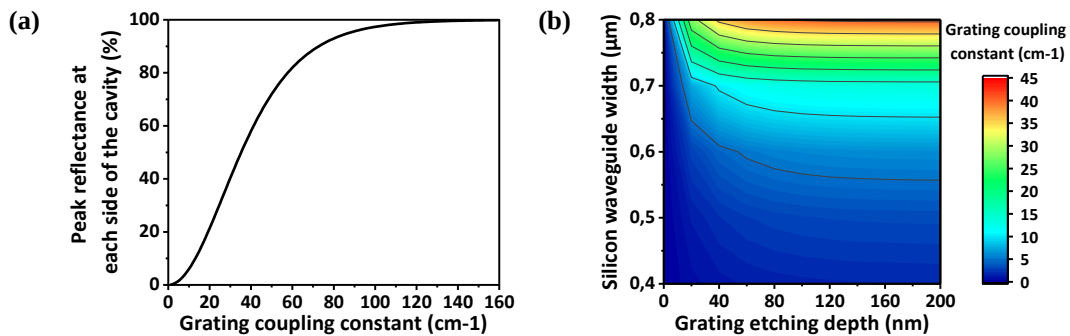


Figure IV-7: (a) Peak reflectance at each side of the cavity versus the grating coupling constant for a grating length of $500/2 = 250\mu\text{m}$. (b) Grating coupling constant of the even supermode versus the etching depth of the grating for different widths of the silicon waveguide below the III-V waveguide. The targeted Bragg wavelength is $1.31\mu\text{m}$.

Since the silicon waveguide width under the III-V waveguide is fixed at $0.8\mu\text{m}$, it imposes a constraint on the adiabatic taper used for coupling, which must start from a larger width than in the DBR case. The shape of the taper used for the DFB is displayed on **Figure IV-8(a)**. In this case, the taper is said to be truncated, but still exhibits an adiabatic character, since it is designed using the methods detailed in **section III-1.2**. The performance of this transition is shown on **Figure IV-8(b)**, where more than 90% of the light is expected to be confined in the silicon waveguide at the end of the transition. The main difference with the taper presented in **section III-1.2** is that the confinement in the MQW along the III-V waveguide is reduced to $\approx 11.7\%$, due to the larger width of the silicon waveguide below.

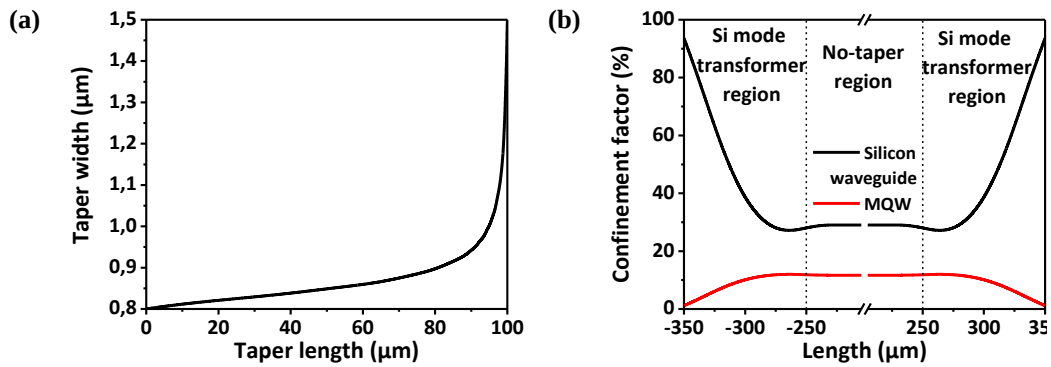


Figure IV-8: (a) Adiabatic taper shape used for the DFB laser. (b) Confinement factor in the MQW and Si waveguide across the whole III-V region.

IV-3. Hybrid III-V on amorphous Si laser fabrication

This section is dedicated to the fabrication of the hybrid III-V on amorphous silicon DFB lasers. Almost all the fabrication processes used for these lasers are the same as those used for the transmitter, and described in **section III-3**. Nevertheless, some differences exist, and are detailed in this section. At first, different possible schemes to integrate the a-Si:H layer are presented, followed by the description of the complete fabrication process, with a focus on the differences with the hybrid transmitter fabrication.

IV-3.1. Amorphous silicon integration schemes

As shown in **section III-3**, the fabrication of a hetero-integrated III-V/Si laser is usually done in three main steps: 1) realization of the waveguides on SOI, 2) bonding of the III-V stack on the patterned SOI and 3) processing of the III-V stack and metallization. Obviously, the a-Si:H layer must be integrated before the bonding of the III-V waveguide.

Two different options were considered to realize the silicon waveguides, both starting from the same thin SOI, and displayed on **Figure IV-9**. In the first one (referred to as option 1), the thin silicon waveguides are first patterned, a SiO_2 cladding is deposited, and then planarized by CMP. Then, a a-Si:H layer is deposited and etched to form the DFB laser cavity and the mode transformers. In the second one (referred to as option 2), a thin oxide layer and the a-Si:H layer are first deposited, followed by a total of four successive etching steps of the two silicon layers to form all the components. Both options end with an oxide layer encapsulation, followed by planarization through CMP, producing the bonding interface for the III-V stack.

Between the two, option 1 is the closest to a full integration case, since the a-Si:H layer is integrated later, and can even be integrated after the formation of active components such as the silicon modulators and germanium photodiodes. However, we decided to proceed with the second option for two reasons. Firstly, once the a-Si:H layer is deposited, the fabrication process is similar to our usual integration process for the laser, and do not need additional developments. Secondly, this option can be considered as a “worst case” for the a-Si:H layer, since it is deposited at the beginning of the process, and will undergo a higher thermal budget than it is supposed to (with the deposition of SiO_2 hard masks at 400°C).

In both options, one can notice the presence of a thin oxide layer between the two silicon layers. The interest of this layer is to act as an etch-stop layer during the a-Si:H layer patterning, in order to protect the

monocrystalline layer. A trade-off on this SiO_2 layer thickness must be respected: it must be thick enough to act as an etch-stop layer, while being as thin as possible to minimize its impact on the silicon waveguide optical properties. In our case, this thickness is fixed at 20nm , but 10nm might have been sufficient.

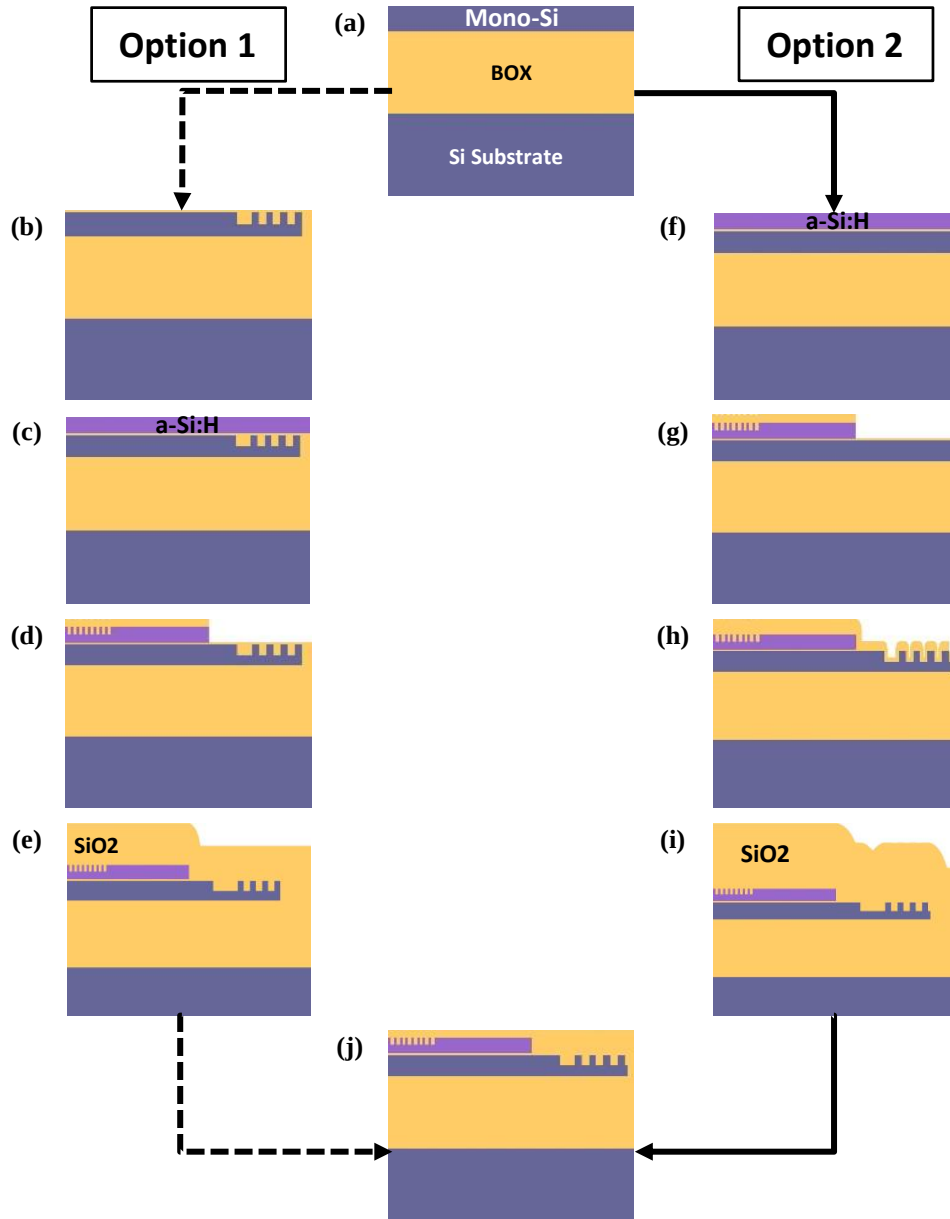


Figure IV-9: Possible integration schemes of the amorphous silicon layer. (a) Starting SOI wafer. (b)-(e) 1st option, where the waveguides are first patterned in the monocrystalline region, followed by the amorphous silicon layer deposition and patterning. (f)-(i) 2nd option, where the amorphous silicon layer is first deposited and patterned, followed by the patterning of the monocrystalline region. (j) Final state, with the planarized SOI wafer ready for bonding.

IV-3.2. Detailed process flow

SOI patterning, encapsulation and planarization [see Figure IV-10(a)-(h)]

The devices have been fabricated on 8" SOI wafers from SOITEC, with a 300-nm -thick silicon layer and a $1\text{-}\mu\text{m}$ -thick buried oxide (BOX). The fabrication starts with the deposition of a 20-nm -thick SiO_2 layer and a 220-nm -thick a-Si:H layer, which is deposited at a temperature of 350°C . The wafers are then processed with four successive steps of 193nm DUV photolithography, and RIE. The first one is a 75-nm -deep etching step, used to define the DFB gratings. A SEM image of the DFB grating is displayed on **Figure IV-11(a)**. Unlike the

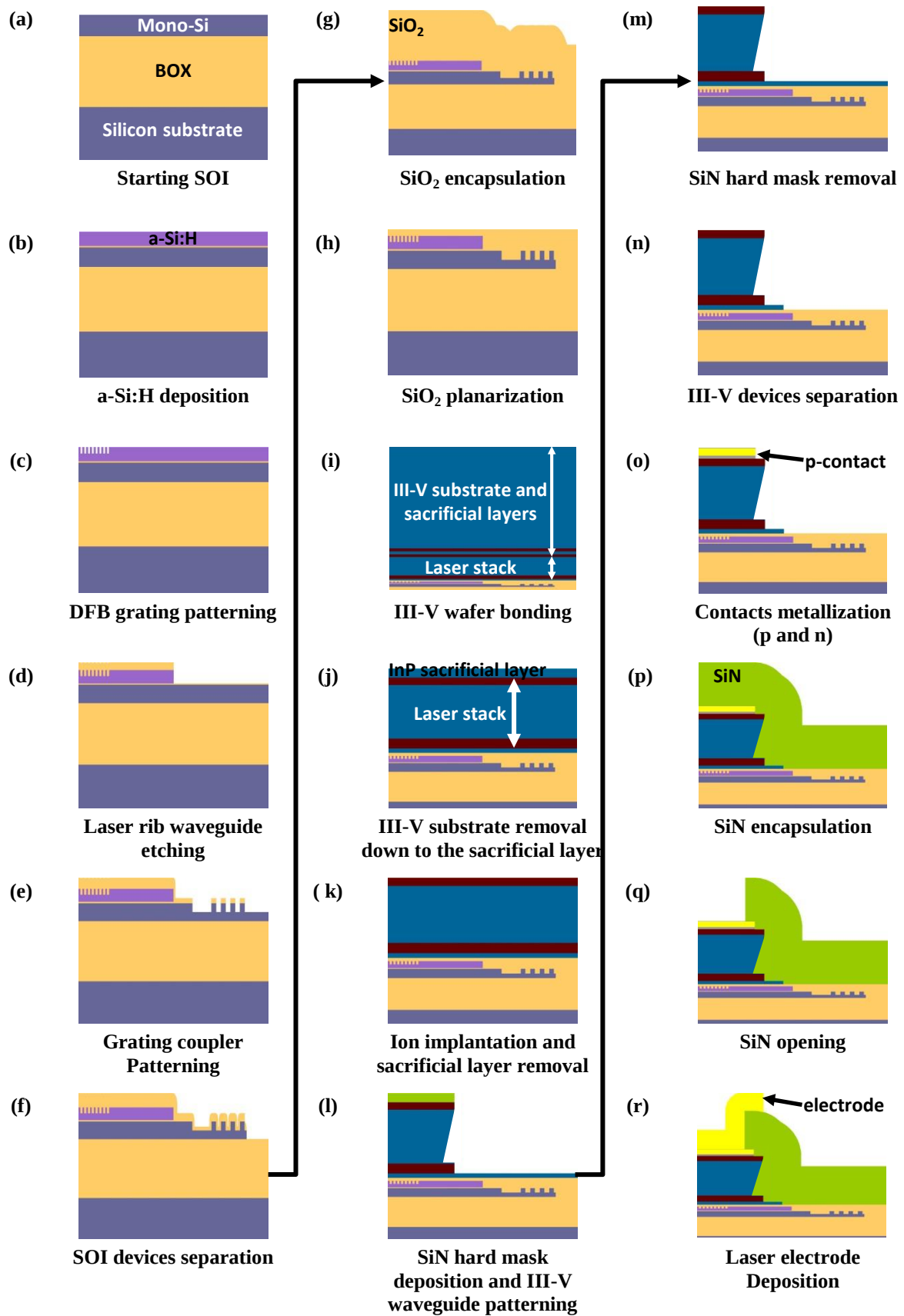


Figure IV-10: Complete fabrication process flow of the hybrid III-V on silicon lasers based on a-Si:H and monocrystalline silicon.

DBR lasers presented in **section III-3**, EBL was not necessary even though the period of the gratings are similar. Indeed, the DFB gratings are defined on a narrower silicon rib waveguide than the DBR gratings, and did not collapse during the lithography. A full-etching of the amorphous silicon layer (stopped by the 20-nm-thick SiO₂ layer) is then realized to form the mode transformers and the transitions from the 500-nm-thick to the 300-nm-thick waveguides. A mode transformer is shown on **Figure IV-11(b)**. The third etching level is a 150-nm-deep etching to form the rib waveguides and the surface grating couplers, while the fourth one is a full-etching of the silicon left (150nm) to separate the devices. A grating coupler can be seen on **Figure IV-11(c)**. Except for the first one, all etching steps used a 100-nm-thick SiO₂ hard mask, deposited at 400°C. The patterned SOI is then encapsulated by 900nm of SiO₂ (also deposited at 400°C), followed by a CMP step, leaving a planar surface for bonding, with approximately 100nm of SiO₂ left above the α -Si:H layer at the center of the wafer.

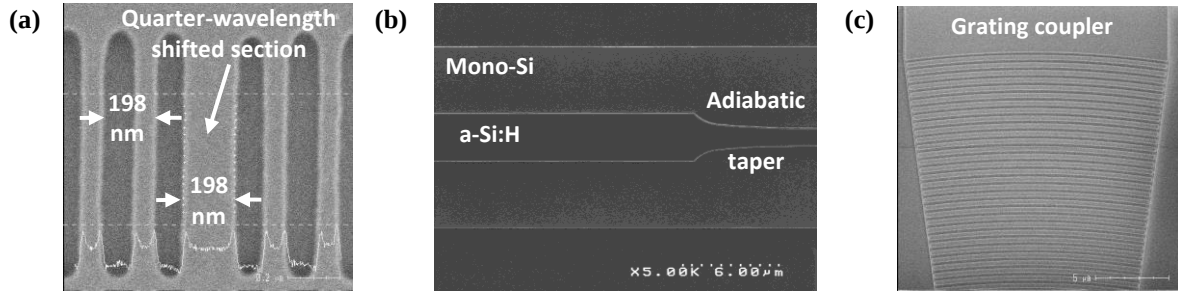


Figure IV-11: SEM views of the (a) center of the DFB grating, (b) adiabatic taper used for the coupling between III-V and silicon waveguide, and (c) grating coupler. The SiO₂ etch-stop and hard-mask layers, have been removed by wet etching.

III-V bonding, substrate removal and hydrogen implantation [see Figure IV-10(i)-(k)]

The composition of the III-V wafers used for the hybrid III-V on amorphous silicon lasers is displayed on **Table IV-1**. These wafers have been ordered to the III-V Lab. The 2.5- μ m-thick epitaxial structure used for the laser is grown from a 2" InP wafer and is close to the one used for the transmitter and shown on **Table III-2**. The main differences are the higher doping concentrations of the p-doped layers, and the thickness of the p-doped InP layer (1.8 μ m instead of 2 μ m).

Table IV-1: III-V epitaxy layer structure used for the integrated III-V on amorphous silicon laser. The grey layers are removed after the bonding, and only the other layers are used to form the active waveguide.

Layer	Material	Photoluminescence peak wavelength (μ m)	Bandgap (eV)	Thickness (nm)	Doping (cm^{-3})
Substrate	InP				
Transition	InP			50	Undoped
Stop-etch	InGaAs			300	Undoped
Sacrificial layer	InP			300	Undoped
p-doped contact	InGaAs	1.65	0.75	200	3×10^{19}
Transition	InGaAsP	1.1	1.13	50	1×10^{19}
p-doped cladding	InP	0.92	1.35	1800	$5 \times 10^{18} \rightarrow 1 \times 10^{18}$
SCH	InGaAsP	1.1	1.13	100	Undoped
Barriers (x7)	InGaAsP	1.1	1.13	10	Undoped
Wells (x8)	InGaAsP	1.28	0.97	8	Undoped
SCH	InGaAsP	1.1	1.13	100	Undoped
n-doped contact	InP	0.92	1.35	110	3×10^{18}
Super-lattice (x2)	InGaAsP	1.1	1.13	7.5	3×10^{18}
Super-lattice (x2)	InP	0.92	1.35	7.5	3×10^{18}
Bonding interface	InP	0.92	1.35	10	Undoped

The bonding and substrate removal processes are the same as those described in **section III-3**. This time, the III-V wafer is bonded at the center of the SOI wafer, where the SiO₂ thickness above the α -Si:H layer is

approximately 100nm, as targeted in the design. The result of the bonding and substrate removal are both displayed on **Figure IV-12**. Except for the top left of the III-V wafer, almost all the surface of the III-V stack is bonded on the SOI wafer.

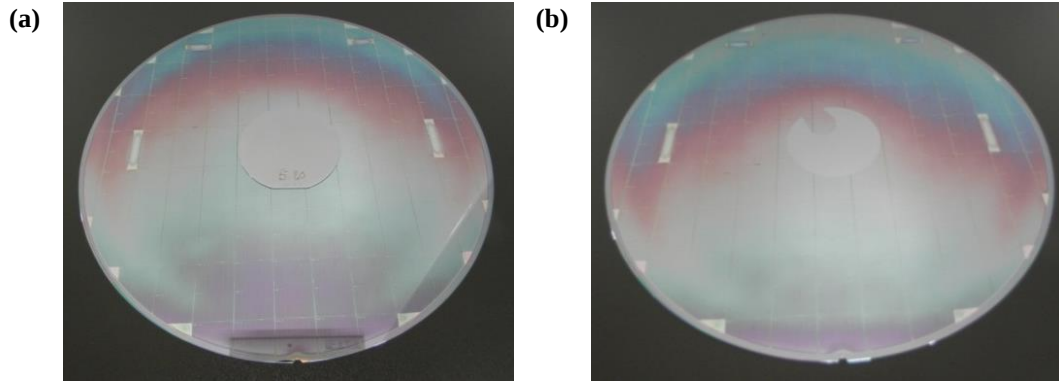


Figure IV-12: (a) Top view of the 8'' SOI wafer after III-V bonding, and (b) after the removal of the substrate.

Unlike for the transmitter case, the sacrificial InP layer is not removed immediately. Actually, an H^+ implantation is used to electrically isolate the edges of the III-V waveguide which is patterned later. The fabrication steps used for the III-V stack implantation are detailed on **Figure IV-13**. A 500-nm-thick SiN hard-mask layer is deposited at 300°C, and patterned by 248nm DUV photolithography and RIE. The thick resist (4μm) used for the SiN etching is also used for the implantation. Three successive implantations are realized with the same dose ($8 \times 10^{13} \text{ cm}^{-2}$) but different energies (220, 230, and 240keV). After resist stripping, the SiN layer is removed by RIE. The *p*-doped contact layer is protected by the sacrificial InP layer during the SiN hard-mask removal. Finally, the sacrificial InP layer is chemically removed by HCl/H₂O wet etching.

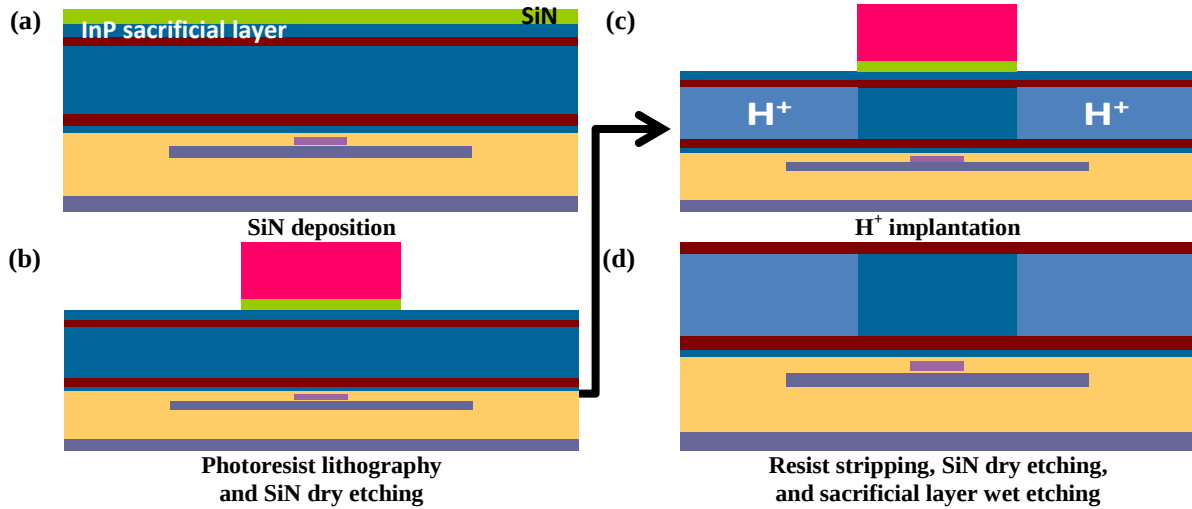


Figure IV-13: H^+ implantation to improve the electrical isolation of the III-V waveguide edges.

III-V waveguide patterning and contact metallization [see Figure IV-10(l)-(o)]

In the case of the transmitter, the SOI wafer was downsized after all sacrificial layers were removed. Here, the SiN hard mask used for the III-V waveguide patterning was deposited (at 300°C) and patterned beforehand, when the SOI wafer was still at the 8'' format, in order to use the lithography equipment from the 8'' cleanroom, which has better performances in terms of alignment. Once the SiN hard mask is patterned, the SOI wafer is downsized from 200 to 75mm, in order to enable subsequent processing in the III-V foundry. The downsized wafer is displayed on **Figure IV-14(a)**. The main consequence is that the III-V waveguide must be patterned before the *p*-contact metallization.

The III-V waveguide is defined using the same combination of dry and wet etching used for the transmitter. Profilometric measurements realised after the p -doped InP wet etching revealed that this layer thickness was closer to $1.6\mu\text{m}$ than the expected $1.8\mu\text{m}$. After the MQW patterning, we estimate that there is a 20nm thickness difference of the n -doped InP contact layer between the center and the edge of the $2''$ stack. Once the MQW have been etched, the SiN hard mask is removed by RIE, and the devices are separated by using a resist mask and Br-based wet etching.

The metallization of the p - and n - contact layers are realized using the metallic stacks and lift-off approach described in **section III-3**. Since it requires an annealing step after deposition, the p -contact metallization is realized first. A 420°C annealing during 10 seconds was thus realized between the p - and n -contacts metallization. The resulting structure is shown on **Figure IV-14(b)**.

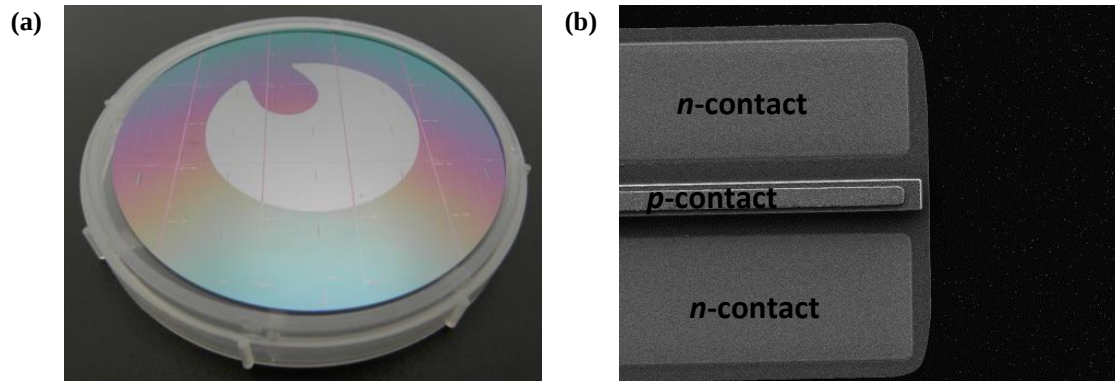


Figure IV-14: (a) Top view of the wafer used for electro-optical characterization (presented in **section IV-4**), after wafer downsizing, and (b) SEM top view of the hybrid III-V on silicon laser after deposition of the n -contact layers.

SiN encapsulation, via opening and electrodes deposition [see **Figure IV-10(p)-(r)**]

After the n -contact metallization, the whole structure is encapsulated in a $2\text{-}\mu\text{m}$ -thick SiN layer deposited at 300°C . Then, the cladding layer is locally opened by RIE to reach the metallic layers. Finally, a $1.3\text{-}\mu\text{m}$ -thick gold layer is deposited and patterned by lift-off to form the electrode, as depicted on **Figure IV-15**.

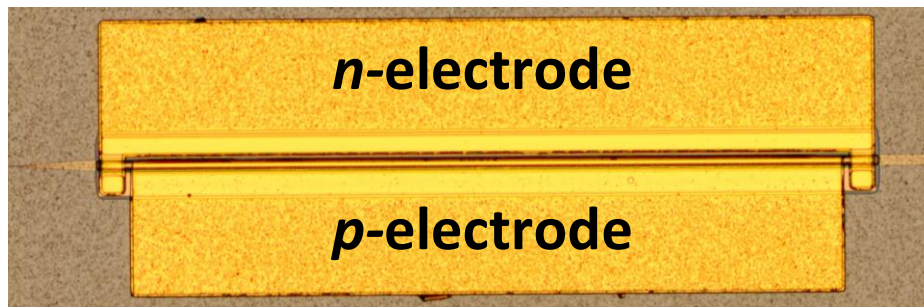


Figure IV-15: Hybrid III-V on amorphous silicon laser after electrodes patterning by lift-off.

During the complete process, several operations necessitating a furnace have been applied on wafers, and thus might possibly increase the absorption material losses of the a-Si:H layer. Among them, the post-bonding annealing and SiN deposition have been realized at 300°C , but should not have a large impact on the a-Si:H layer, according to the study presented in **section IV-1.2**. In the same way, all the SiO_2 depositions have been realized at 400°C , but are not expected to have a large effect, since they provide hydrogen atoms during their deposition. The total time spent at 400°C during all the SiO_2 depositions was estimated at ≈ 35 minutes. Finally, only the p -contact annealing at 420°C during 10 seconds is expected to possibly influence the losses of the a-Si:H layer.

IV-4. III-V on amorphous Si laser characterization

This section presents the characterization of the hybrid III-V on amorphous silicon DFB lasers. All the fabricated devices demonstrated laser operation when electrically pumped in CW regime, with level of optical power over 1mW collected in the SMF for most of them. Therefore, the possible material losses coming from the amorphous silicon layer do not prevent laser operation. Furthermore, the best of these lasers, which is thoroughly presented in this section, showed performances at the state-of-the-art for hybrid III-V on silicon DFB lasers. Unlike the modulators and transmitters characterizations which have been done at room temperature, the tests of hybrid III-V on amorphous silicon lasers have been realized on a support with a controllable temperature, fixed at 25°C. Pictures of the laser under test are presented on **Figure IV-16**.

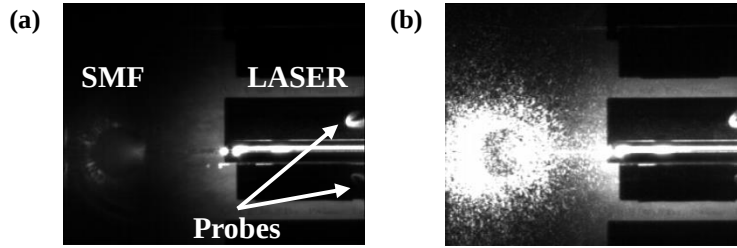


Figure IV-16: Laser under test observed through an infrared camera, (a) below threshold and (b) above threshold.

The characterization starts with the waveguide-to-fiber grating coupler, measured as a stand-alone structure, using the same method described in the previous chapters. The transmission spectrum of the measured grating coupler is shown on **Figure IV-17**. This spectrum is used as a reference to estimate the power coupled at the output of the laser in the thin monocrystalline silicon waveguide. Its maximum of transmission is approximately at -5.2dB . Around 1280nm (which is the peak wavelength of the laser shown in the following measurements), the grating coupler loss is approximately at -5.5dB , which is the value used to estimate the power in the silicon waveguide.

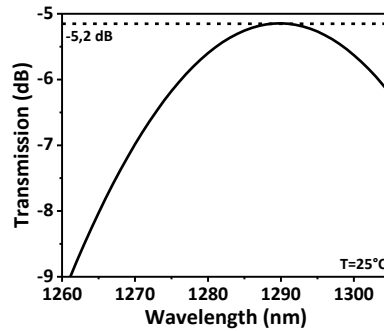


Figure IV-17: Optical loss of an individual grating coupler (fitted measurement).

The static performances of the laser are evaluated using the set-up described on **Figure III-32**, with a grating coupler located at one end of the DFB to collect the emitted power in a SMF. The $P(I_d)$ characteristic of the laser is displayed on **Figure IV-18(a)**. Its threshold current at 25°C is 50mA, corresponding to a threshold current density of 1.43kA/cm^2 , and the maximum power coupled in the SMF is approximately 4.7mW at 186mA. Based on the measured losses of the grating coupler at the laser peak wavelength (-5.5dB), and by assuming that the output power is the same at both sides of the DFB, the maximum power coupled in the thin monocrystalline silicon waveguide is estimated to at 33.6mW. Using the same estimation, the slope $\Delta P/\Delta I_d$ above threshold is 0.26W/A , resulting in a differential quantum efficiency of 26.5% according to **Eq. (I.10)**. While it is not apparent on the $P(I_d)$ characteristic of the laser, it must be noted that for an injected current above 144mA, the laser suffers from a mode competition, as shown later on its spectral measurements.

The $V(I_d)$ characteristic of the laser is shown on **Figure IV-18(b)**. The series resistance is 5Ω , which is similar to the resistance obtained for the lasers shown in **section III-4.1**. The turn-on voltage is 1.24V, which

gives a turn-on electrical power of $62mW$. The wall-plug efficiency is also evaluated using Eq. (I.11), and displayed on **Figure IV-18(c)**. The WPE reaches a maximum of 10% for an injected current of $140mA$.

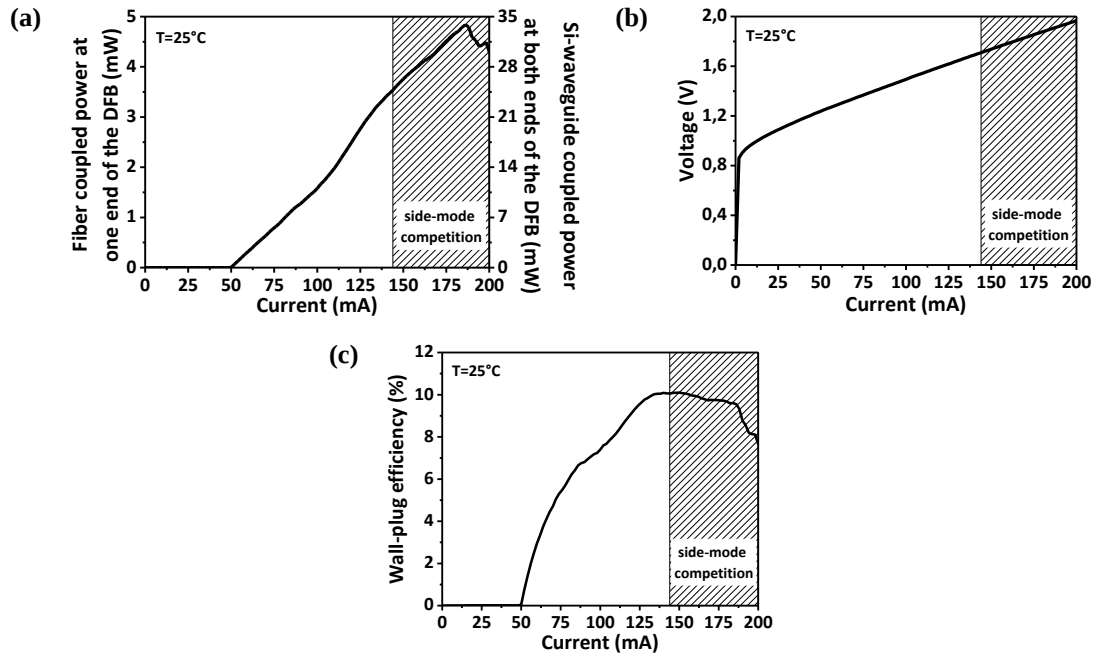


Figure IV-18: (a) Optical power measured at the grating coupler output (left axis) and estimated in the silicon waveguide (right axis), (b) voltage between the lasers electrodes, and (c) Wall-plug efficiency, all versus pumping current.

The lasing spectrum is shown on **Figure IV-19(a)** as a function of the pumping current. Above threshold, the laser spectrum is typical of a QWS DFB, reaching single-mode operation with a maximum SMSR evaluated at $51dB$ for a $120mA$ injected current, as displayed on **Figure IV-19(b)**. However, as shown in **Figure IV-19(c)**, even if the SMSR is above $45dB$ up to $144mA$, it falls down to $15-20dB$ afterwards.

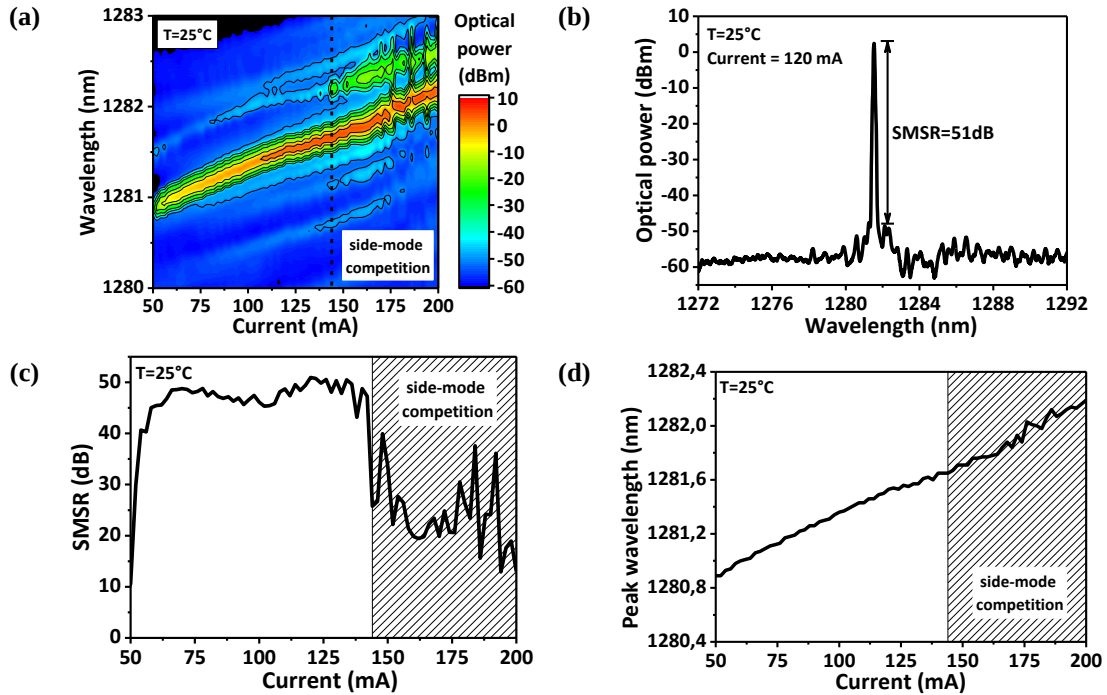


Figure IV-19: (a) Laser optical spectrum as a function of the pumping current. (b) Optical spectrum at a fixed pumping current ($120mA$). (c) SMSR and (d) peak wavelength as a function of the pumping current.

As it can be seen on **Figure IV-19(a)**, this SMSR drop is due to a mode competition between the lasing mode and one of the side-mode close to the edge of the grating bandwidth (the one with the longer wavelength). We assume this competition is induced by the longitudinal spatial-hole burning effect, which is typical of QWS DFB lasers at a high injection current. Above threshold, the optical field is amplified in the phase-shifted section and the injected free carriers will move to this region. The local variation of free carriers induces a non-uniform distribution of both refractive index and gain along the laser cavity. Therefore, when the current increases, the gain margin between the lasing mode and the non-lasing side-modes is reduced, until multi-mode oscillation occurs. This effect is the main issue of QWS DFB lasers, which limit their power operation in the single-mode regime [167].

Nevertheless, as it can be seen on **Figure IV-19(d)**, the laser stays mode-hop free, even with the side-mode competition. When the injected current is increased from 50 to 200mA, the wavelength shifts from 1280.9nm to 1282.2nm. It can be noted that while the carriers injected in the junction modify the refractive index of the gain region – which in turn modifies the effective index of the even supermode, and thus the Bragg condition – their effect should shift the wavelength towards lower wavelengths. This shift towards higher wavelengths is actually due to the variation of refractive index of the gain region with the temperature due to the self-heating of the laser, which counteract the effect of the injected carriers. Moreover, as explained in **section IV-2.2**, the Bragg wavelength for this period was supposed to be close to 1.31 μ m. According to **Eq. (III.9)**, the difference in peak wavelength can be explained by a smaller average effective index for grating than expected in our simulation (≈ 3.23 instead of 3.3).

The $P(I_d)$ characteristic of the laser has also been evaluated by increasing the temperature of the support from 15°C to 55°C. When the temperature increases, the effects of mechanisms such as Auger recombination and carrier leakage outside the active region before stimulated recombination also increase, which in turn increase the threshold current and reduce the differential quantum efficiency. These effects can be seen on **Figure IV-20**. The limit temperature for laser operation is slightly below 55°C, which is close to the state-of-the-art of hybrid III-V on silicon lasers based on InGaAsP MQW.

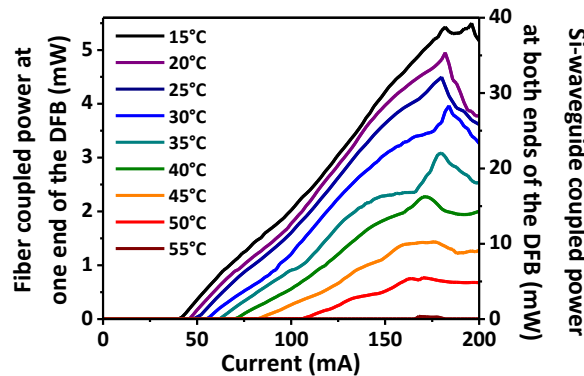


Figure IV-20: Optical power measured at the grating coupler output (left axis) and estimated in the silicon waveguide (right axis) versus pumping current, at different temperatures.

IV-5. Conclusions on the use of amorphous silicon

In this chapter, a hybrid III-V on silicon DFB laser based on a bi-layer made of amorphous and monocrystalline silicon has been effectively demonstrated. The deposition of this amorphous silicon layer has been identified as a good solution to effectively couple the light from a thick III-V waveguide to a thin silicon waveguide, without starting from a thicker SOI layer. Indeed, amorphous silicon can provide low material absorption loss and can be integrated in the later steps of fabrication due to its low-deposition temperature (below 400°C).

The demonstrated DFB laser also showed excellent performances which are close to the state-of-the-art lasers of the same type. Thus, its threshold current (50mA), current density (1.43kA/cm²) and maximum SMSR (51dB) are on the same level as the DFB lasers presented in **Table I-1**. Its maximum power estimated in the waveguide is also on the same level, even if limited in the single-mode regime (at 24mW) most probably due to

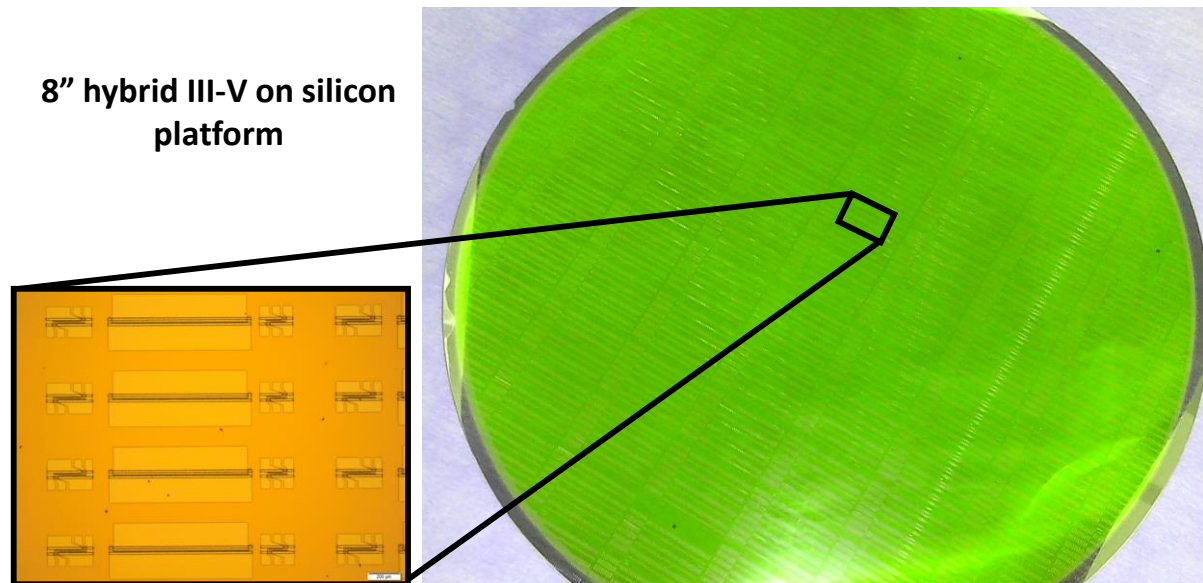
the spatial-hole burning effect. Furthermore, its differential quantum efficiency (26.5%) and Wall-Plug efficiency (10%) are similar to the best demonstrated for DFB lasers up to now [60]. The laser operates until temperatures slightly below 55°C, which is reasonable, since the lasers presented on **Table I-1** could operate up to 55-60°C [50], [60].

This demonstration shows the possibility to use amorphous silicon to couple light generated from a hetero-integrated III-V on silicon laser into a thin silicon layer, such as the ones used by silicon photonics fabrication platforms. The thermal budget applied (several SiO₂ depositions at 400°C and an annealing at 420°C for 10 seconds) during the laser fabrication did not prevent the laser from operating with great performances. However, the maximum thermal budget which can be applied on the laser without degrading its performances is not known, and precautions should still be taken if the amorphous silicon layer is deposited before the metal layers deposition of the silicon photonic circuit, which thermal budget is around 400°C, for larger annealing times than the one used for the laser *p*-contact annealing. Nevertheless, if the amorphous silicon layer is deposited after these steps, our demonstration showed that the laser processing steps are not an issue.

Finally, it can be noted that the introduction of this amorphous silicon layer can open new perspectives in the design of silicon photonics components, which might take advantage of a locally increased silicon thickness. For instance, this idea has already been demonstrated some time ago for silicon grating couplers, which have shown enhanced performances in terms of coupling efficiency [168].

Chapter V: Conclusions and outlook

This chapter is the last of the document, and concludes the entire study. For starters, the studies, results and conclusions obtained during the PhD thesis, and presented in the previous chapters, are summarized. Then the next steps for the transmitter integration are detailed, as well as future directions to investigate in order to further improve the components performances. Finally, this document will be concluded with a more general view on what can be expected for the future of silicon photonics.



V-1. Work overview	114
V-1.1. Results synthesis.....	114
V-1.2. Conclusion	115
V-2. Next steps for the transmitter integration	115
V-2.1. Complete fabrication process on 8" or 12" wafers	115
V-2.2. Electro-optical integration	117
V-3. Further improve the components performances.....	117
V-3.1. Laser thermal performances.....	117
V-3.2. Pulse-amplitude modulation (PAM) format.....	119
V-4. Future of silicon photonics	120

V-1. Work overview

The aim of this PhD thesis was to conceive and produce a high-speed silicon photonic transmitter with an integrated laser source, targeting applications such as data communication in the next generation of data-centers based on single-mode fiber transmission. Indeed, current silicon photonic transmitters generally rely on externally laser sources, which are limited by non-negligible coupling optical losses and tight tolerances toward spatial misalignments, resulting in increased cost and complexity of the complete silicon photonic circuit packaging. On the other hand, a transmitter with a wafer-scale integrated laser source which would not suffer from these issues is a major objective for silicon photonics. Since there are currently no viable monolithically integrated laser sources on silicon, the most promising solution to form a wafer-level laser source is the heterogeneous integration of III-V materials via bonding. Rather than directly modulate the hybrid III-V on laser source, we decided to integrate it with a modulator, in order to form a high-speed III-V on silicon integrated transmitter. As for its specifications, they were inspired from the 100GBASE-LR4 norm. Thus, the transmitters were supposed to operate at four different wavelengths (with a 4.5nm wavelength spacing), centred on $1.3\mu\text{m}$, with a data transmission rate of 25Gb/s per wavelength, and for distances up to 10km .

V-1.1. Results synthesis

The first step was to choose and conceive the modulators used for the transmitters, amongst the different existing structures for light modulation. While an interesting solution would have been to use hybrid III-V on silicon electro-absorption modulators, using the same III-V stack for both components requires a trade-off on their respective performances. Thus, it would have been necessary to use multiple die bonding techniques requiring a precise alignment. It would also have been necessary to increase the complexity of the III-V processing, while one of the main interests of silicon photonics is to transfer the maximum of functions in the silicon, to benefit from its mature fabrication process. Therefore, silicon Mach-Zehnder modulators based on carrier-depletion in a p - n junction have been chosen for their inherent robustness toward fabrication variations, as well as their high-speed capabilities even with their relatively weak modulation efficiency. These modulators have been realized on a thin SOI layer (300nm), and were conceived by taking into account the limitations inherent to the future laser co-integration (with only one metal level available). The design of both p - n junctions and travelling-wave electrodes were detailed. Stand-alone modulators (without the integrated lasers) have also been fabricated, and their characterization validated the design, showing that the modulators were perfectly suited for 25Gb/s operation. A direct comparison with the state-of-the-art for these structures was not obvious, since most of them were tested for data transmission rates above 40Gb/s , while we were limited by the eye-diagram testing set-up. Nevertheless, by looking at the specific characteristics of the modulators and using eye diagrams simulations based on the measurements, it could be seen that the modulators had honourable performances in terms of speed and efficiency compared to the state-of-the-art, given their design limitations. The main issue concerning the modulators performances were the relatively large level of optical losses of the passive circuitry, leading to high optical powers needed from the laser source. Additional simulations have shown that large improvements are possible by improving the performances of the passive components. Finally, solutions have been proposed to further improve the electro-optical bandwidths of the modulators, by using impedance adaptation, slow-wave electrodes, or reducing the junction capacitance.

The modulators being conceived, they were co-integrated with a hybrid III-V on silicon laser, in order to form the transmitter. A distributed Bragg reflector laser structure was chosen for its wavelength tuning capabilities, in order to be compatible with the targeted wavelength spacing (4.5nm). The design of the lasers, based on InGaAsP multiple quantum-well which maximum gain was centred on $1.3\mu\text{m}$, was presented, with a focus on the adiabatic tapers used for III-V to silicon optical coupling, and the Bragg mirrors forming the laser cavity. These components were defined in a 500-nm -thick silicon waveguide, necessary for light coupling. A transition taper was used to couple the light signal into a 300-nm -thick silicon waveguide, where all the other components (silicon modulators and waveguide-to-fiber grating couplers) were defined. The process flow was thoroughly presented, as well as the characterization of the fabricated transmitters and their individual components. Using these transmitters, data transmission at 25Gb/s was effectively demonstrated, for distances up to 10km , and at two different wavelengths. Modulation at other wavelengths has not been realized, but could have been done by tuning the laser peak wavelength, which was not possible during eye diagrams measurements.

due to equipment limitations. While the demonstration was indeed successful, the transmitters also had several limitations on their performances, which sources have been identified thanks to the characterization of the individual components forming the transmitters. Some of these limitations (low laser output optical power, single-wavelength regime, modulator bandwidth, heater efficiency) were due to incidents during the fabrication process or to the components design, and can be corrected relatively easily. Nevertheless, the most important issue of the transmitter was its reduced output power, which was obviously due to the low laser output optical power, but also to the optical propagation losses which were larger than for the stand-alone silicon modulator measurements. This increase of optical losses might be attributed to the additional etching processes, necessary to form the 300-nm-thick silicon components from a 500-nm-thick SOI starting layer.

Another objective was to improve the integration of the laser source with the silicon photonic circuit. The issue of needing a silicon waveguide thicker than the SOI layers used in standard silicon fabrication platforms for optical coupling was tackled. Rather than relying on complex processing of the III-V waveguide or modifying the III-V epitaxial structure, we proposed a solution based on the deposition of an amorphous silicon layer to locally, in order to locally increase the thickness of the silicon waveguide. Low-loss amorphous silicon can be deposited at low-temperatures (below 400°C), and be integrated in the later steps of fabrication, without disturbing the fabrication of the standard components based on a thin SOI layer. A quarter-wavelength shifted distributed feed-back structure was used for the demonstration of a hybrid III-V on silicon laser based on a bi-layer made of amorphous and monocrystalline silicon. The design of this specific laser structure was detailed, as well as its process flow. Finally, the characterization of the best device showed excellent performances, which are at the state-of-the-art of hybrid III-V on silicon distributed feed-back lasers.

V-1.2. Conclusion

Although the various limitations of the fabricated devices, this work demonstrates that current silicon photonics technology, associated with heterogeneous integration on silicon, enables to build systems which are compatible with the main constraints of data-communication over kilometric distances, without relying on optical sources external to the system and integrated as discrete components. Moreover, the various issues either coming from the design or the fabrication process have been identified, and can be corrected to propose a new set of transmitters with improved performances. Finally, it has also been demonstrated that amorphous silicon was a viable solution to couple light generated from hetero-integrated III-V on silicon laser into thin silicon layers, thus solving one of the laser integration issues presented in **section I-3.4**.

V-2. Next steps for the transmitter integration

The integrated transmitter presented in this study is based on hybrid III-V on silicon lasers, integrated by wafer bonding. During the fabrication, the 8" SOI wafers were downsized to 3" to finish the laser processing and realize the contacts metallization. This type of process flow is sufficient for prototyping, but is not suitable for large scale fabrication. Not only a large surface of both SOI and III-V wafers are wasted, but the current laser integration process is realized with non-standard fabrication methods (such as lift-off) and materials for the contacts, and prevents the use of standard multi-level metallization as explained in **section I-3.4**. In this section, the developments necessary to solve these integration issues, as well as the solutions currently investigated at CEA-LETI, are briefly presented. Those developments are essential for the future of integrated III-V on silicon transmitters, or else the laser source will remain external in silicon photonic transmitters. A more standard fabrication process would also permit the integration of the driving circuits at the wafer scale, using the three dimensional stacking of integrated circuits, which is another key advantage of silicon photonics.

V-2.1. Complete fabrication process on 8" or 12" wafers

In the process flow used for the transmitter, and presented in **section III-3.1**, wafer bonding was used to heterogeneously integrate III-V material on the SOI wafer. As explained in **section I-3.2**, this method is not suited for integration on 8" or even 12" SOI wafers. Therefore, the die bonding approach is currently developed, by focusing on how to improve the bonding yield and throughput.

For now, the SOI wafers are downsized a few steps after the III-V bonding. This downsizing is necessary to realize the III-V etching process and contact metallization in another cleanroom suited for III-V processing, due

to material contamination issues. Obviously, the processing of III-V materials is not considered as “standard”, since these materials are not yet widespread in microelectronics. However, this might not be the case in the future, in order to form the next-generation of CMOS transistors. Therefore, III-V RIE processes are currently developed in the 8” cleanroom to form the active waveguides. As for the contact metallization, another PhD student (Elodie Ghegin) is currently investigating the use of metals considered as more “CMOS-friendly” than the current gold-based metallization. This would represent a large step forward for the lasers integration. Nevertheless, it can be noted that even if the contacts are not yet “CMOS-friendly”, a complete processing of stand-alone lasers on 8” SOI wafers without any downsizing has been recently demonstrated at CEA-LETI and is illustrated in the front page of this chapter [169].

However, even if the lasers can be integrated on 8” or 12” wafers as stand-alone components, the last integration issue listed in **section I-3.4** remains. Indeed, the III-V stack used for the laser is still over $2\mu\text{m}$ thick, and prevents the use of conventional multi-level metallization. Even if the hybrid III-V on silicon transmitter presented in **chapter III** was realized using non-planar metallization processes generally used for the lasers (with one metal level and lift-off methods), this approach is only valid for prototyping. Moreover, the silicon modulator RF performances would also be improved by using several metal levels to form the electrodes. While it might be possible to realize planar multi-level metallization by setting the first metal level above the laser, this approach would ask to re-develop the metallization architecture and process, which are generally standardized. An alternative solution is currently investigated by another PhD student (Jocelyn Durel), and is referred as back-side laser integration. In this integration case, all the silicon photonic components are first formed with all their standard metallization stacks. A carrier silicon wafer is then bonded on top of the SOI wafer, before flipping the assembly. The substrate and BOX of the SOI wafer can then be removed, before bonding the III-V materials and patterning the laser. A schematic view of the back-side laser with other components is presented on **Figure V-1**.

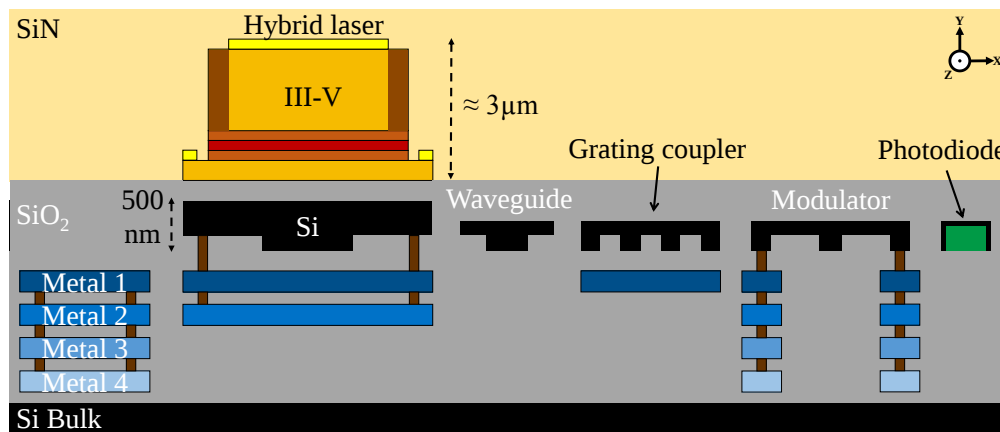


Figure V-1: Schematic view of the back-side laser integration. Image from [170].

This back-side laser integration does not only solve the multi-level metallization issue, but also provides several interesting features. First of all, the III-V processing steps are moved at the end of the fabrication process, and can thus be processed in the far back-end areas of the fabrication platforms, which should simplify the integration of CMOS-unfriendly materials. It can also be noted that the amorphous silicon approach used to locally increase the silicon waveguide thickness can also be applied here, and integrated after the metallization of the silicon active components, thus reducing its thermal budget. Finally, by using a metallic mirror deposited above the grating coupler (using the first metal layer), light can be coupled into a single-mode fiber located at the back of the silicon photonic circuit, which can be useful if an electrical integrated circuit is stacked above the silicon photonic circuit, as discussed in the next section. Waveguide-to-fiber grating couplers using metallic [171] or polycrystalline silicon [172] back-side mirrors have already been used to improve the coupling efficiency of the grating coupler. In the case of the back-side laser integration scheme, grating couplers with a metallic mirror as depicted on **Figure V-1** have recently been demonstrated with the same performance that in a standard front-side integration and better results are expected for grating couplers with a specific design for back-side integration [170].

V-2.2. Electro-optical integration

If hybrid III-V on silicon photonic circuits can have access to a planar metallization, they will also be able to use one of the key advantages of silicon photonics circuits: the three dimensional stacking of integrated circuits at the wafer-level. An example is displayed on **Figure V-2**. While the electrical circuit driving the silicon photonic circuit can be integrated by wire bonding as in [173], this method can only be applied once the wafers have been diced, and is not suited for large-scale production. Moreover, three dimensional stacking also provides low parasitic capacitances and inductances, which do not limit the devices RF performances [174]. It can also be noted that with the back-side integration, light extraction would be simplified, since the optical signal would be coupled on the opposite side of the electrical integrated circuit.

Stand-alone silicon modulators have already been driven by these integrated circuits, in travelling-wave [174], [175] and segmented designs [176], where several short active sections ($\approx 500\text{-}\mu\text{m}$ -long) of the Mach-Zehnder modulator are driven by different drivers. In this case, due to the small capacitance of each section, travelling-wave electrodes are not necessary to reach large electro-optical bandwidths. While adding more complexity to the electrical integrated circuit, segmented devices might also permit to increase further the maximum data rate of the modulators (up to 100Gb/s), with an appropriate design of the driver output impedance to reduce the RC-limitation [177].

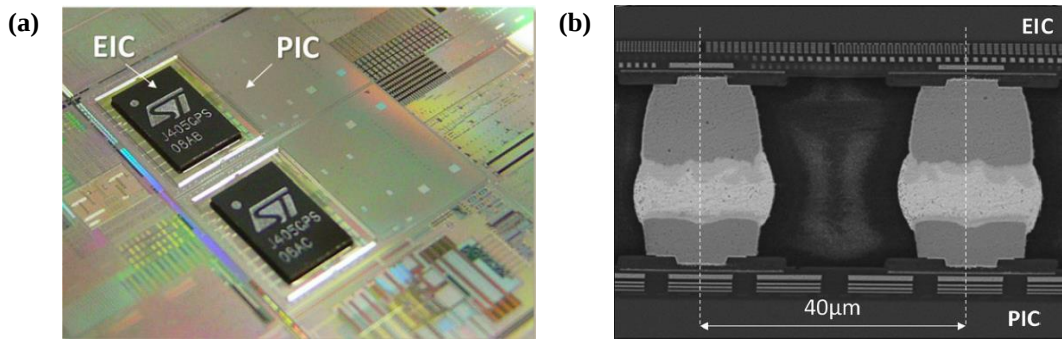


Figure V-2: (a) Top view of a three-dimensional assembly of an electrical integrated circuit (EIC) and a photonic integrated circuit (PIC). (b) Transversal SEM view of copper pillars used for die-to-wafer flip-chip bonding. Images from [11].

V-3. Further improve the components performances

During the study of the different components, several solutions have already been proposed to solve the short-term issues limiting the performances of the transmitters, without large modifications on the design or the fabrication process. The aim of this section is to look at more “advanced” solutions to improve even further the performances of the components on aspects which have not been extensively been treated in this work, such as the maximum temperature of operation of the lasers and more advanced modulation formats than On-Off Keying, used to increase the data transmission rate of the components.

V-3.1. Laser thermal performances

While the subject was not directly tackled in the manuscript, it was shown that the performances of the demonstrated hybrid III-V on silicon lasers are degraded with an increasing temperature, and even stopped to operate around 55°C . This is an issue for the laser sources, since they will not necessarily operate in environments with a controlled temperature. Moreover, if an integrated circuit is bonded on top of the laser, it will also heat the overall circuit when operating [11]. This temperature limitation is mainly due to the materials used for the multiple quantum-wells (InGaAsP) and the BOX layer, which prevents the heat from flowing in the silicon substrate, due the poor thermal conductivity of silicon dioxide.

Compared to other material systems, InGaAsP lasers are more sensitive to the electron leakage current from the active region, which increases with the temperature, and reduces the laser performances. The current leakage mainly depends on the conduction band offset between the active and cladding regions, which controls the confinement of the electrons in the active region. Indeed, electrons having a lighter effective mass than holes,

they require a tighter confinement [3]. A solution would be to use the InGaAlAs material system for the laser, which was briefly introduced in **section II-1.3** for the electro-absorption modulators. This material system enables a larger conduction band offset than in InGaAsP systems, which in turn reduces the leakage current density. Therefore, InGaAlAs lasers with a specifically designed epitaxial stack can operate at larger temperatures, as in [44], where continuous-wave regime is reached up to 105°C. A possible issue with InGaAlAs lasers might be their fabrication process, which seems to be less reliable than for InGaAsP lasers, thus making them more prone to failure [2].

In order to solve the BOX issue, a recently proposed solution was to connect the *p*- and *n*- electrical pads of the laser to the silicon substrate, by locally etching both the SOI and BOX layers before the metal layers deposition [178]. This solution is illustrated on **Figure V-3(a)**. Hybrid III-V on silicon microring lasers based on these metal shunts have shown a large improvement of their maximum temperature of operation (+30°C), compared to designs without metal shunts. Lasers with metals shunts also had higher output powers, and could operate at higher currents before the differential efficiency roll-over, which indicates a reduction of the self-heating effects. However, the operating temperature improvement might also be due to the small dimensions of the microrings, which have a larger access resistance than other types of laser cavities. Indeed, another proposed solution was to use polycrystalline silicon shunts between the laser and the silicon substrate, formed before the III-V bonding [179]. In this case, the lasers (800- μm -long Fabry P  rot cavities) with a thermal shunt did not show an improvement on their maximum temperature of operation, but could also be operated at higher currents before experiencing the differential efficiency roll-over. Therefore, the metallic thermal shunt design should be tested on other types of laser cavities to see its influence on their maximum temperature of operation. Nevertheless, the thermal shunt method remains a good solution to reduce the lasers self-heating issues and improve their performances.

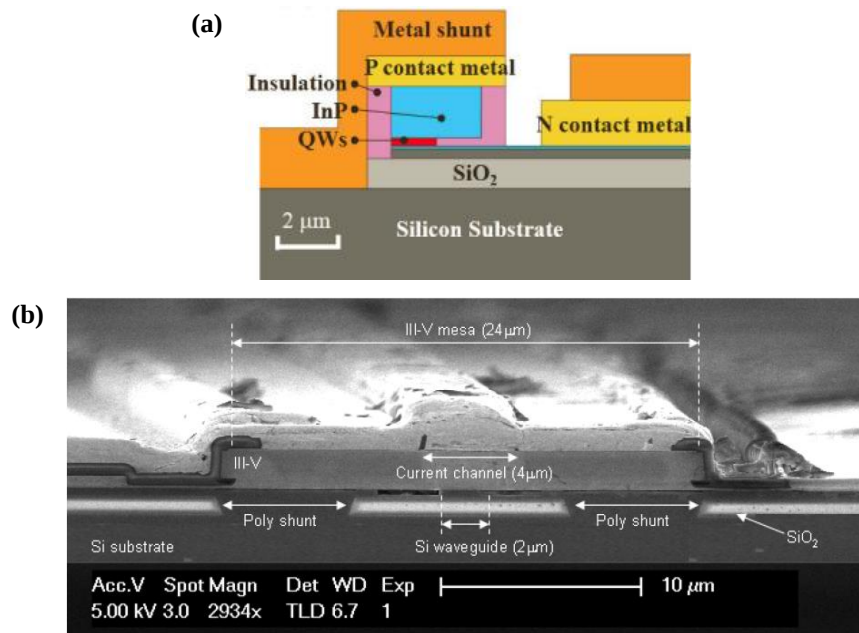


Figure V-3: (a) Solutions proposed in [178] to improve the thermal behavior of the hybrid III-V on silicon lasers, by connecting the electrical pads to the silicon substrate. (b) Other solution based on poly-silicon thermal shunts to the silicon substrate proposed in [179].

Finally, a promising solution to improve even further the thermal behaviour of the lasers would be to use quantum-dots (QD) structures. Indeed, due to the three dimensional confinement of carriers, QD lasers are expected to operate at higher temperatures than multiple quantum-well lasers, but to also provide lower threshold current densities. QD lasers are already used for the III-V monolithically integrated on silicon sources presented in **section I-2.3**, since they are more tolerant toward dislocation defects. Indeed, only a limited number of QDs will effectively be affected by these defects [25]. An example of InAs/GaAs QD structure is shown on **Figure V-4(a)**. Thus, QD lasers grown on silicon have already demonstrated lasing operation up to 119°C [180]. In the same way, externally integrated QD laser diodes have also been used to demonstrate transmitters operating up to

125°C [33]. Even more recently, the first hybrid III-V on silicon QD laser integrated by direct bonding was demonstrated [181]. Transversal views of the structure are presented on **Figure V-4(b)** and **(c)**. Continuous-wave operation was demonstrated up to 100°C, with a Fabry-Pérot structure, which performances are comparable in terms of output power to the ones based on multiple quantum-wells presented in **Table I-1**, and with threshold current densities as low as 271 A/cm².

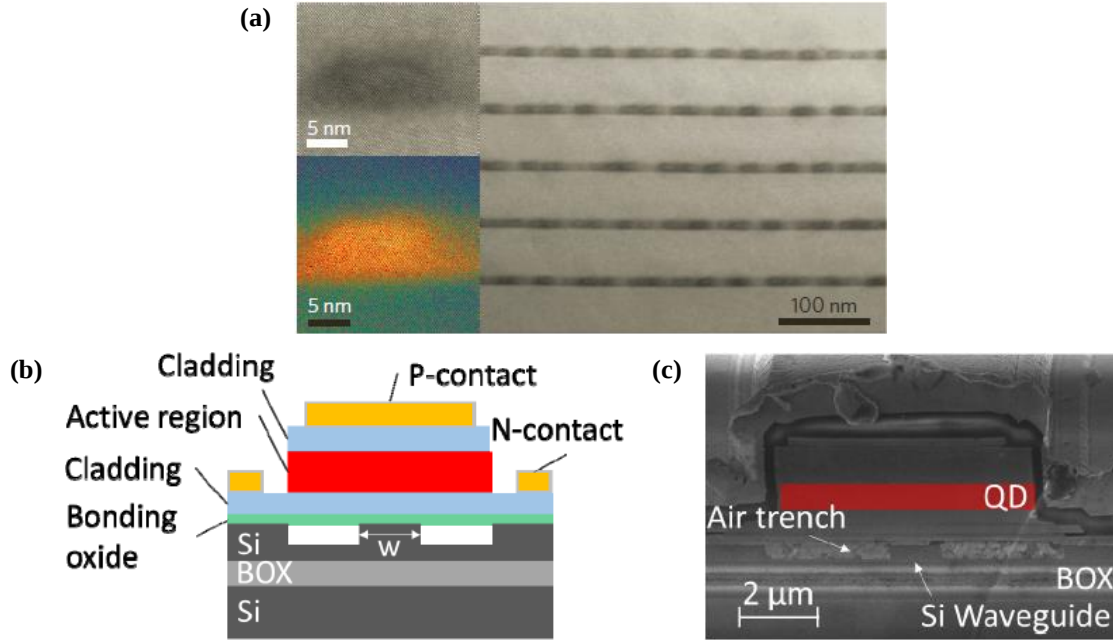


Figure V-4: (a) Transmission electron microscopy (TEM) image of an InAs/GaAs QD structure, with insets of an individual QD [25]. (b) Schematic and (c) SEM transversal views of a hybrid III-V on silicon QD laser [181].

V-3.2. Pulse-amplitude modulation (PAM) format

As explained at the beginning of this document, the volume of information which circulates across the world will only continue to increase. While the standard treated in this document was limited to 100Gb/s Ethernet applications, the next 400Gb/s Ethernet standard is already being defined [182]. In order to reach such transmission rates two solutions are possible: either increase the number of wavelengths, or improve further the maximum data rate of each wavelength. For now, both solutions are considered, with 8 wavelengths modulated at 50Gb/s, or 4 wavelengths modulated at 100Gb/s. Data rates of 50Gb/s have already been demonstrated for silicon modulators, therefore the first solution should be reachable for silicon photonic transmitters. However, using 8 wavelengths simultaneously will also increase the overall system complexity, and this solution might not be retained. Therefore, solutions to increase the bandwidth of the existing devices must be found.

In the modulator study presented in **chapter II**, only the On-Off Keying modulation format was presented. However, other modulation formats exist and can improve the transmitted data rate without requiring larger electro-optical bandwidths. For instance, coherent modulation schemes such as quadrature phase-shift keying (QPSK) and quadrature amplitude modulation (QAM) are currently used for long-distance transmission systems. However, they are generally not considered as a solution for transmission in the data-centers, due to the complexity and cost of the dedicated receivers. In the recent years, another modulation scheme referred as pulse-amplitude modulation (PAM) format has gained interest for optical communications. Instead of encoding the information on 2 optical power levels as in On-Off Keying, intermediary levels can be used to increase the number of bits transmitted, without modifying the bit duration. The most common is the PAM-4 modulation format, where 4 optical power levels are used to encode information. Thus, at a given bit duration, the bit rate for a PAM-4 modulation scheme will be twice the one for an On-Off Keying format. The PAM-4 signal can be generated in the electrical domain with an electrical digital-to-analog converter (DAC), which will generate the levels and sent them to a silicon modulator [126]. Another solution is to use two Mach-Zehnder modulators with different lengths, but receiving the same electrical levels, which will generate the 4 power levels directly in the

optical domain. This approach is generally preferred since it simplifies the electronic drivers design. Using this kind of structure, 100Gb/s data transmission has been successfully demonstrated [183], and the resulting eye diagram is displayed on **Figure V-5**. Therefore, while being more complex than the On-Off Keying format, the PAM format would represent a good solution to increase the modulator maximum bit rate, if it would be impossible to sufficiently improve its electro-optical bandwidth in order to satisfy the next standards of transmission rates.

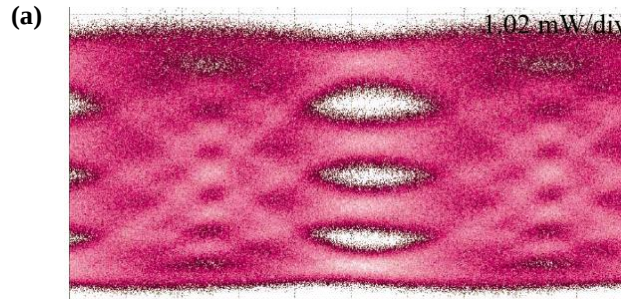


Figure V-5: 100Gb/s eye diagram for PAM-4 modulation obtained with silicon Mach-Zehnder modulators. Image from [183].

V-4. Future of silicon photonics

In 2006, it was guessed that within 5 years, fast and cost-effective optical interconnections between computer chips would be available using silicon photonics [184]. The potential benefits coming from such a large market were one of the reasons why silicon photonics has known such a fast-paced development in the last decade. Silicon photonics was supposed to revolutionize the microelectronics industry.

Ten years later, despite the large progress made and the excellent performances reached by silicon photonics components, the expected revolution has yet to happen. There are still no commercially available systems based on optical interconnect between processors, and silicon photonics has been struggling to find its place against the existing photonic systems in the short and long-reach domains. Since 2012, doubts have been casted on the real capacity of silicon photonics to satisfy industries expectations. This situation can be explained by the optimistic previsions made on silicon photonics developments. Indeed, it was first envisioned that electronic and photonic component would be monolithically co-integrated on the same wafers and provide dense interconnects for communications between several chips for instance. However, as explained in **section I-2.2**, the SOI and BOX layers suited thickness are not the same for each type of components, which is why they have been separated in the first place. Finding a trade-off in their integration would also be a hindrance towards the development of advanced CMOS nodes. It was also expected that monolithically integrated laser sources on silicon would soon be available, in order to further increase the interconnection density with several compact sources. Another issue is the thermal management of the silicon photonic circuits which are extremely sensitive to temperature variations, limiting their efficiency in computing environments [185].

Nevertheless, silicon photonics is more and more seen as a cost-efficient solution compared to other photonic systems, and is starting to impose itself in the data-centers. It might even overcome the dominant multi-mode VCSELs due to their transmission distance limitations. Commercial solutions based on silicon photonics have been present for some years in active optical cables, and are now proposed as highly efficient optical modules. With the progresses made on integrated laser sources and thermal management of the silicon photonic circuits, very-short reach optical interconnections may at least be within reach in the next few years. Thus, the development of the hybrid III-V on silicon microdisks and microrings lasers presented in **section I-3.2** are still ongoing [178], [186], and might become the compact sources used for intra-chip communications. Finally, while this entire document was dedicated to data communications, it must also be remembered that silicon photonics is not limited to these applications. For instance, it can also emerge as a solution in other applications, such as information processing, display, or sensing. Silicon photonics have a large unveiled potential, which will still be investigated for some years.

This work was supported by the French national program 'Programme d'Investissements d'Avenir', IRT Nanoelec, Grant ANR-10-AIRT-05.

List of figures

Figure I-1: (a) Attenuation coefficient, and (b) dispersion coefficient dependence on the wavelength for silica-glass fibers. Images from [5].	3
Figure I-2: Example of photonic optical link for the 100GBASE-LR4 standard.	4
Figure I-3: Interband transitions between the conduction (top) and valence (bottom) bands, in which the energy is supplied by or given to a photon: (a) spontaneous emission, (b) absorption, (c) stimulated emission. Filled circles represent the electrons, and open circles the holes. Images from [3].	5
Figure I-4: Energy band diagrams and major carrier transition processes in a direct band structure such as InP (left) and in an indirect band structure such as silicon (right). Image from [15].	6
Figure I-5: Schematic view of a Raman scattering laser. Image from [16].	6
Figure I-6: Band structures of (a) bulk Ge, (b) tensile strained intrinsic Ge, and (c) tensile strained n-doped Ge. Images from [18].	7
Figure I-7: Schematic view of InP directly grown on silicon laser (a) with a thick III-V buffer [25] and (b) without buffer [26].	8
Figure I-8: Examples of silicon photonic circuits with external laser diodes (referred as LD) integrated to the silicon substrate. (a) LD in a hermetic micro-package [27], (b) Arrayed LD chip [28].	8
Figure I-9: Coupling schemes between an edge-emitting laser diode and a silicon photonic circuit.	9
Figure I-10: Schematic view of an optical oscillator. Image from [5].	10
Figure I-11: Schematic view of a laser cavity, with an active region (with a length L_a) and a passive region (with a length L_p). The gain region is restricted to the dashed area of the active region. Feedback is provided by two mirrors with respective complex reflection coefficients r_1 and r_2 . Image from [3].	10
Figure I-12: (a) Schematic view of the different elements needed to form a laser for single wavelength operation. (b) Spectral characteristics of each element. Images from [3].	11
Figure I-13: Typical plots used to characterize the laser diode. (a) Optical output power and (b) voltage versus drive current. (c) Laser output spectrum at a fixed driving current.	12
Figure I-14: (a) 3" III-V wafer bonded on an 8" SOI wafer. (b) 104 III-V dies (with dimensions of $5\text{mm} \times 5\text{mm}$) bonded on an 8" SOI wafer. Second image taken from [70].	13
Figure I-15: Schematic views of a hybrid III-V over silicon structure. (a) Three-dimensional view [15], and (b) Cross-sectional view, with the optical mode distributed between the active and passive waveguides.	14
Figure I-16: Limit cases of waveguide coupling where: (a) the even/odd supermode is mostly confined in waveguide 1/2, (b) both supermodes are equally shared between the two waveguides (phase matching), and (c) the even/odd supermode is mostly confined in waveguide 2/1.	15
Figure I-17: Schematic view of the power transfer between the active and passive waveguides using adiabatic coupling.	16
Figure I-18: (a) III-V waveguide metal bonded to the silicon substrate, edge coupled to a SOI waveguide [68]. (b) Active waveguide bonded on a silicon substrate, with re-grown contact regions after bonding for lateral injection [82].	16
Figure I-19: Schematic view of a Fabry-Pérot cavity with cleaved facets in the silicon waveguide. Image from [62].	17
Figure I-20: Top microscopic view of a racetrack cavity. Image from [46].	17
Figure I-21: (a) DBR laser with shallowly-etched gratings [57]. (b) Ring-assisted DBR laser [52].	18
Figure I-22: Schematic view of a DFB laser with a quarter-wavelength shift. Image from [43].	18
Figure I-23: Schematic views of hybrid III-V on silicon (a) microdisk [75], and (b) microring [15] lasers.	19

Figure I-24: Schematic view of a silicon photonic circuit comprising different kinds of passive and active devices, and illustrating the issues of laser integration in terms of SOI layer thickness and standard multi-leveled metallizations. This co-integration is not possible in these conditions.	21
Figure I-25: Schematic views of high-speed transmitters co-integrating hybrid III-V on silicon lasers and silicon modulators, with : (a) four lasers with a fixed wavelength [85], and (b) a single tunable laser [86]......	22
Figure I-26: Hybrid III-V on silicon integrated transmitter, where the III-V is not only used to form the laser, but all active components (modulators and photodiodes). Schematic views of the transmitter (a) after bonding, (b) after substrate removal, and (c) after etching and metallization processes. Images from [88].	23
Figure II-1: General principle of optical intensity modulation.	26
Figure II-2: Eye diagrams measurements at 25Gb/s of the same modulator, in the same driving conditions, but with different level of power measured by the detector: (a) low optical power (without output amplification), (b) high optical power (with output amplification).	28
Figure II-3: Schematic (a) Top and (b) cross-sectional views of a hybrid III-V/Si EAM. Images from [97].	30
Figure II-4: Schematic cross-sectional views of (a) a Ge-based FKE modulator [103] and (b) a Ge/SiGe QSCE modulator [104].	31
Figure II-5: Schematic cross-sectional views of modulators based on (a) carrier injection [110], (b) carrier accumulation [112], and (c) carrier depletion [127].	32
Figure II-6: (a) Top schematic view [132], and (b)-(c) transmission spectra [131] of a ring modulator.	33
Figure II-7: (a) Top microscopic view of a silicon MZM. (b) Schematic view of the p - n junction embedded in its arms. (c) Transmission spectrum of an asymmetric MZM. (d) Transmission spectrum of a symmetric MZM. Images (a)-(c) are from [123], image (d) is from [134].	34
Figure II-8: Transversal schematic view of the silicon MZMs doped waveguide.	36
Figure II-9: (a) Fundamental TE mode propagating in the silicon waveguide. (b) Additional optical losses coming from the highly doped regions of the silicon MZMs.	36
Figure II-10: Absolute net doping distribution in the waveguide, simulated with (a) an uniform doping concentration , and (b) a realistic doping profile. (c) Comparative horizontal cuts of both profiles.	37
Figure II-11: Phase shift at a -2.5V bias and doping-related losses coming from doping at a 0V bias, for (a,c) uniform doping concentrations, and (b,d) realistic doping profiles.	38
Figure II-12: Phase shift at -2.5V bias versus doping related losses at 0V bias for realistic doping profiles. Each point represents a different combination of acceptor (Na) and donor (Nd) atoms concentrations.	38
Figure II-13: (a) Linear phase shift, (b) doping-related losses, and (c) junction capacitance versus reverse-voltage bias applied on the chosen junction for the modulator.	39
Figure II-14: (a) Transmission line model. (b) RLGC model of the modulator active region.	40
Figure II-15: Schematic transversal view of the silicon MZM, with its definitive dimensions.	41
Figure II-16: Simulations of the linear (a) resistance and (b) inductance of the electrodes, and of the linear (c) conductance and (d) capacitance between the electrodes. Lines: case for a HR substrate ($750\Omega.cm$). Dashed lines: case for a SR substrate ($10\Omega.cm$).	41
Figure II-17: Equivalent analytic circuit of the active region.	42
Figure II-18: Simulations of the (a) real and (b) imaginary part of the characteristic impedance, and of the (c) RF losses and (d) RF effective index. Lines: case for a HR substrate ($750\Omega.cm$). Dashed lines: case for a SR substrate ($10\Omega.cm$).	43
Figure II-19: Simulations of the S-parameters magnitudes for the different lengths of the active region: 2-mm-long (a) and (b), 4-mm-long (c) and (d) , 6-mm-long (e) and (f). Lines: case for a HR substrate ($750\Omega.cm$). Dashed lines: case for a SR substrate ($10\Omega.cm$).	44
Figure II-20: Simulations of the modulation depth for the different lengths of the active region: (a) 2-mm-long, (b) 4-mm-long, and (c) 6-mm-long. Lines: case for a HR substrate ($750\Omega.cm$). Dashed lines: case for a SR substrate ($10\Omega.cm$).	45
Figure II-21: Layout view of the complete MZM structure.	46
Figure II-22: Complete fabrication process flow of the silicon MZMs.	47

Figure II-23: Implantation steps for Si-doping	48
Figure II-24: Evolution of the (a) phase shift at $-2.5V$ bias and (b) doping-related losses at $0V$ bias, versus lithographic misalignment of the doped p and n regions from their targeted positions.	48
Figure II-25: Silicon waveguide patterning.	49
Figure II-26: SEM views of the (a) grating coupler, (b) MMI, and silicon waveguide forming the (c) MZMs arms.	49
Figure II-27: Silicide formation.	50
Figure II-28: SEM views of the MZMs silicided regions.	50
Figure II-29: Metal deposition and electrodes patterning.	50
Figure II-30: Set-up for optical passive measurements.	51
Figure II-31: (a) Example of raw measurement on a single die of the power received by the power-meter (input power: $0dBm$). (b) Optical loss of the individual grating couplers on each die (fitted).	52
Figure II-32: Set-up for static electro-optical measurements.	52
Figure II-33: Transmission spectra of the device under different bias from input fiber to output fiber, for	53
Figure II-34: Set-up for S-parameter measurements.	54
Figure II-35: Simulations and measurements of the S-parameters magnitudes for the different lengths of the active region: (a) and (b) 2-mm-long, (c) and (d) 4-mm-long, (e) and (f) 6-mm-long.	54
Figure II-36: Simulations and measurements of the (a) real and (b) imaginary part of the characteristic impedance, and of the (c) RF losses and (d) RF effective index.	55
Figure II-37: Set-up for small-signal electro-optical measurements.	56
Figure II-38: Simulations and measurements of the modulation depth for the different lengths of the active region: (a) 2-mm-long, (b) 4-mm-long, and (c) 6-mm-long.	56
Figure II-39: Simulation of the transmission spectra of the device under different bias from input fiber to output fiber, for (a) 2-mm-long, and (b) 4-mm-long silicon MZM.	57
Figure II-40: Typical form of the electrical signals applied on each MZMs arm, in a dual-drive push-pull operation at $25Gb/s$	57
Figure II-41: Simulated transmission at $0V$ and ER when the MZM is operated in dual-drive push-pull, with $-2.5V$ applied on each arm for (a) a 2-mm-long active region and (b) a 4-mm-long active region.	57
Figure II-42: Simulated transmission at $0V$ and OMA when the MZM is operated in dual-drive push-pull, with $-2.5V$ applied on each arm for (a) a 2-mm-long active region and (b) a 4-mm-long active region.	58
Figure II-43: (a) and (b) MZM under high-speed testing. (c) Set-up for large-signal electro-optical measurements.	58
Figure II-44: Eye diagrams at $25Gb/s$ of the electrical signal send on the RF probes: (a) without de-emphasis, (b) with de-emphasis.	59
Figure II-45: Measured and simulated eye diagrams at $25Gb/s$ of the (a,b) 2-mm-long and (c,d) 4-mm-long silicon MZM, in dual drive configuration with $2.5V_{pp}$ applied on each arm, each biased at $-1.25V$	59
Figure II-46: Simulated eye diagrams at $40Gb/s$ of the (a) 2-mm-long and (d) 4-mm-long silicon MZM, both in dual drive configuration with $2.5V_{pp}$ applied on each arm, each biased at $-1.25V$	61
Figure II-47: Simulated eye diagrams at $50Gb/s$ of the (a) 2-mm-long and (d) 4-mm-long silicon MZM, both in dual drive configuration with $2.5V_{pp}$ applied on each arm, each biased at $-1.25V$	61
Figure II-48: Calculated OMA at $25Gb/s$ for (a) the fabricated MZMs, and (b) MZMs with the same $p-n$ junction and TWE electrodes, but improved passive components. In both cases, the additional optical losses coming from a 10-km-long propagation in a SMF are not taken into account.	61
Figure III-1: (a) Transversal, (b) longitudinal, and (c) top schematic views of the laser.	64
Figure III-2: Effective index of the simulated fundamental TE modes of the III-V waveguide and a silicon slab waveguide (10- μm -wide) versus its thickness, at $1.31\mu m$. Both waveguides are cladded in SiO_2	65
Figure III-3: (a) Effective index of the simulated TE modes of the isolated III-V and silicon waveguide versus the width of the silicon rib section, at $1.31\mu m$. Both waveguides are cladded in SiO_2	67

Figure III-4: (a) Effective index of the simulated even and odd supermodes of the hybrid III-V and silicon waveguide, and (b) confinement factor of each supermode in the silicon waveguide and in the MQW (without the barriers), versus the width of the silicon waveguide, at $1.31\mu\text{m}$.	67
Figure III-5: Dependence of the silicon waveguide width on the γ parameter.	68
Figure III-6: (a) Normalized shape of the adiabatic taper. (b) Power confinement factor in the silicon waveguide at the end of the III-V region versus the input and output widths of the adiabatic taper.	68
Figure III-7: Electric field distribution of the fundamental even mode at the center (a), and end (b) of the gain region. (c) Longitudinal view showing the power transfer from the end of the III-V region to the Si waveguide. (d) Confinement factor in the MQW and Si waveguide across the whole III-V region.	69
Figure III-8: Power confinement factor in the silicon waveguide at the end of the III-V region versus	69
Figure III-9: Longitudinal (a) and transversal (b) schematic views of the Bragg reflectors.	70
Figure III-10: Reflectance spectra for two gratings, with different etching depths, but the same product $\kappa g L g \approx 5.9$, thus the same peak reflectance. The targeted Bragg wavelength is 1310nm for each grating. For $g = 10\text{nm}$ (red line), $\kappa g \approx 79\text{cm}^{-1}$, and $L g = 750\mu\text{m}$.	71
Figure III-11: (a) FWHM and (b) peak reflectance versus the etching depth and grating length of silicon Bragg gratings. The targeted Bragg wavelength is 1310nm for each grating.	71
Figure III-12: Bragg wavelength dependency on (a) grating period for a 10nm etching depth, and (b) etching depth for a 196nm grating period.	72
Figure III-13: (a) Spectra and (b) peak reflectance versus grating length for several etching depths. The targeted Bragg wavelength in the calculations is 1310nm .	73
Figure III-14: Front (black line) and back (red line) mirrors reflectance spectrums, for an etching depth of 10nm , with (a) a grating period of 195nm , and (b) a grating period of 197nm .	73
Figure III-15: (a) Top view of the transition. (b)-(d) Electric field distributions of the fundamental TE mode across the transition.	74
Figure III-16: Schematic transversal view of the silicon MZM, with its electrodes modified due to the co-integration with the hybrid III-V on silicon laser.	75
Figure III-17: Simulations of the linear (a) resistance and (b) inductance of the electrodes, and of the linear (c) conductance and (d) capacitance between the electrodes of the silicon modulator. All simulations are realized with a HR substrate ($750\Omega\cdot\text{cm}$). Lines: previous electrodes described in section II-2.2 . Dashed lines: New electrodes due to the laser co-integration.	75
Figure III-18: Simulations of the modulation depth for the different lengths of the active region: (a) 2-mm -long, (b) 4-mm -long, and (c) 6-mm -long. All simulations are realized with a HR substrate ($750\Omega\cdot\text{cm}$). Lines: previous electrodes described in section II-2.2 . Dashed lines: New electrodes due to the laser co-integration.	76
Figure III-19: Layout view of the complete transmitter co-integrating the hybrid III-V on silicon DBR laser and the silicon MZM (2-mm -long).	76
Figure III-20: SEM views of the (a) end transition from thick to thin silicon waveguides, (b) Bragg reflector, and (c) adiabatic taper used for the coupling between III-V and silicon waveguide.	77
Figure III-21: Complete fabrication process flow of the hybrid III-V on silicon transmitters.	78
Figure III-22: (a) Top view of the $8''$ SOI wafer after III-V bonding, and (b) after the removal of the substrate and the sacrificial layers. (c) Top view of the wafer used for electro-optical characterization (presented in section III-4), after wafer downsizing.	80
Figure III-23: Deposition and patterning of the p-contact metallic layers.	81
Figure III-24: Top (a) optical microscope and (b) SEM views of p-contact metallic layers after deposition and lift-off.	81
Figure III-25: III-V waveguide patterning.	82
Figure III-26: (a) Patterned SiN hard mask, covering the p-contact metal layers. Top optical microscope views of the lasers after separation, focused on (b) the p-contact and (c) the n-doped InP layer.	82
Figure III-27: SEM side view of the hybrid III-V on silicon laser after deposition of the n-contact layers.	83
Figure III-28: Metallization of the silicon modulator.	83

Figure III-29: Hybrid III-V on silicon laser (a) after SiN opening, and (b) after electrodes lift-off.	83
Figure III-30: Top view of the complete transmitter, with SEM views of the different components.	84
Figure III-31: Optical loss of an individual grating coupler (fitted measurement).	85
Figure III-32: Set-up for the laser electro-optical measurements.	85
Figure III-33: (a) Optical power measured at the grating coupler output (left axis) and estimated in the silicon waveguide (right axis) and (b) voltage between the lasers electrodes, both versus pumping current.	86
Figure III-34: Optical spectrums for (a) DBR1 at 100mA, (b) DBR2 at 100mA, (c) DBR1 at 120mA, and (d) DBR2 at 130mA.	86
Figure III-35: Phase shift versus applied voltage, simulations and measurements for a 2-mm-long MZM.	87
Figure III-36: Measurements of the S-parameters magnitudes for the different lengths of the active region: (a)-(b) 2-mm-long, (c)-(d) 4-mm-long, and (e)-(f) 6-mm-long. The purple and blue lines are the same as the ones presented on Figure II-35	88
Figure III-37: Simulations and measurements of the (a) real and (b) imaginary part of the characteristic impedance, and of the (c) RF effective index and (d) RF losses.	88
Figure III-38: Laser voltage and optical power collected in the fiber versus laser pumping current for both (a) transmitter 1 and (b) transmitter 2. For the displayed values of current, the lasers operate in single-mode regime.	89
Figure III-39: (a) Broadened and (b) narrowed spectrum of transmitter 1 for a fixed laser wavelength, with different static phase difference between the MZM arms. The laser drive current is fixed at 100mA.	90
Figure III-40: Spectrum for a fixed MZM phase difference, with a tuned laser wavelength (a) for transmitter 1, with a laser drive current fixed at 100mA, and (b) for transmitter 2, with a laser drive current fixed at 150mA. P_h is the total electrical heating power used for both the front and back mirrors.	90
Figure III-41: (a) and (b) Pictures of the transmitter under high-speed testing. (c) Set-up for large-signal electro-optical measurements.	91
Figure III-42: 25Gb/s eye diagrams of the hybrid III-V on silicon transmitters using (a) 4-mm-long MZM at 1303.5nm (transmitter 1) and (b) 2-mm-long silicon MZM at 1315.8nm (transmitter 2), both in dual drive configuration with 2.5Vpp applied on each arm.	92
Figure III-43: 25Gb/s eye diagram of transmitter 1 in the same condition that in Figure III-42 (a) , but using the BOA instead of the external photodiode.	92
Figure III-44: 25Gb/s eye diagram of (a) transmitter 1 and (b) transmitter 2, in the same conditions that in Figure III-42 , but with a 10-km-long SMF.	93
Figure IV-1: (a) Schematic view of a inverted double-taper architecture, and (b) scanning electron microscope image of its III-V tip [60]. (c) Schematic view of a same direction double-taper architecture and (d) scanning electron microscope image of its III-V tip [153].	97
Figure IV-2: (a) Schematic view and (b) scanning electron microscope image of the III-V waveguide cross-section used in [154].	97
Figure IV-3: Schematic view of the atomic structure of (a) monocrystalline silicon and (b) hydrogenated amorphous silicon. Images from [159].	99
Figure IV-4: (a) Schematic cross-section of the laser based on a-Si:H and monocrystalline silicon. (b) Simulated confinement factor of the even supermode in the rib and slab sections of the silicon waveguide.	99
Figure IV-5: Schematic transversal view and spectra of (a) a standard DFB laser and (b) a quarter-wavelength shifted DFB laser. The entire grating is distributed along the active gain material. δ is the same detuning parameter from the Bragg wavelength (λ) defined in Eq. (III.7) . Images from [3].	100
Figure IV-6: Schematic longitudinal view of the laser based on a-Si:H and monocrystalline silicon layers.	101
Figure IV-7: (a) Peak reflectance at each side of the cavity versus the grating coupling constant for a grating length of $500/2 = 250\mu m$. (b) Grating coupling constant of the even supermode versus the etching depth of the grating for different widths of the silicon waveguide below the III-V waveguide. The targeted Bragg wavelength is $1.31\mu m$	102
Figure IV-8: (a) Adiabatic taper shape used for the DFB laser. (b) Confinement factor in the MQW and Si waveguide across the whole III-V region.	103

Figure IV-9: Possible integration schemes of the amorphous silicon layer. (a) Starting SOI wafer. (b)-(e) 1st option, where the waveguides are first patterned in the monocrystalline region, followed by the amorphous silicon layer deposition and patterning. (f)-(i) 2nd option, where the amorphous silicon layer is first deposited and patterned, followed by the patterning of the monocrystalline region. (j) Final state, with the planarized SOI wafer ready for bonding.	104
Figure IV-10: Complete fabrication process flow of the hybrid III-V on silicon lasers based on a-Si:H and monocrystalline silicon.....	105
Figure IV-11: SEM views of the (a) center of the DFB grating, (b) adiabatic taper used for the coupling between III-V and silicon waveguide, and (c) grating coupler. The SiO ₂ etch-stop and hard-mask layers, have been removed by wet etching.	106
Figure IV-12: (a) Top view of the 8" SOI wafer after III-V bonding, and (b) after the removal of the substrate.	107
Figure IV-13: H ⁺ implantation to improve the electrical isolation of the III-V waveguide edges.....	107
Figure IV-14: (a) Top view of the wafer used for electro-optical characterization (presented in section IV-4), after wafer downsizing, and (b) SEM top view of the hybrid III-V on silicon laser after deposition of the <i>n</i> -contact layers.	108
Figure IV-15: Hybrid III-V on amorphous silicon laser after electrodes patterning by lift-off.	108
Figure IV-16: Laser under test observed through an infrared camera, (a) below threshold and (b) above threshold.	109
Figure IV-17: Optical loss of an individual grating coupler (fitted measurement).....	109
Figure IV-18: (a) Optical power measured at the grating coupler output (left axis) and estimated in the silicon waveguide (right axis), (b) voltage between the lasers electrodes, and (c) Wall-plug efficiency, all versus pumping current.	110
Figure IV-19: (a) Laser optical spectrum as a function of the pumping current. (b) Optical spectrum at a fixed pumping current (120mA). (c) SMSR and (d) peak wavelength as a function of the pumping current.	110
Figure IV-20: Optical power measured at the grating coupler output (left axis) and estimated in the silicon waveguide (right axis) versus pumping current, at different temperatures.....	111
Figure V-1: Schematic view of the back-side laser integration. Image from [170].	116
Figure V-2: (a) Top view of a three-dimensional assembly of an electrical integrated circuit (EIC) and a photonic integrated circuit (PIC). (b) Transversal SEM view of copper pillars used for die-to-wafer flip-chip bonding. Images from [11].	117
Figure V-3: (a) Solutions proposed in [178] to improve the thermal behavior of the hybrid III-V on silicon lasers, by connecting the electrical pads to the silicon substrate. (b) Other solution based on poly-silicon thermal shunts to the silicon substrate proposed in [179].	118
Figure V-4: (a) Transmission electron microscopy (TEM) image of an InAs/GaAs QD structure, with insets of an individual QD [25]. (b) Schematic and (c) SEM transversal views of a hybrid III-Von silicon QD laser [181].	119
Figure V-5: 100Gb/s eye diagram for PAM-4 modulation obtained with silicon Mach-Zehnder modulators. Image from [183].	120

List of tables

Table I-1: State-of-the art of the existing heterogeneously integrated III-V on silicon laser structures.	20
Table I-2: Targeted performances for the transmitter.	24
Table II-1: Doping conditions for a uniform concentration of $1 \times 10^{17} \text{cm}^{-3}$ in a 300-nm-thick silicon waveguide.	37
Table II-2: Simulated S-parameters performances with a HR substrate.....	44
Table II-3: Simulated electro-optical bandwidths with a HR substrate.	45
Table II-4: Doping conditions used for the modulator.....	48
Table II-5: Measured optical losses for each component, and in the complete MZM structures.....	52
Table II-6: Simulated and measured $S_{21} - 6\text{dB}$ bandwidths.	55
Table II-7: Simulated and measured electro-optical bandwidths.....	56
Table II-8: Summarized performances of the MZMs from input SMF to output SMF.....	59
Table II-9: State-of-the art of the silicon MZMs based on lateral p - n junctions, with different driving schemes.....	60
Table II-10: Expected performances of the MZMs from input SMF to output SMF, with the same junction and TWE, but improved performances for the grating couplers and passive silicon waveguides.....	62
Table III-1: Bragg reflectors characteristics.	73
Table III-2: III-V epitaxy layer structure used for the integrated III-V on silicon transmitter.....	79
Table III-3: Measured $S_{21} - 6\text{dB}$ bandwidths for the stand-alone modulators measured in section II-4.3 and the modulators co-integrated with the hybrid III-V on silicon lasers.....	89
Table IV-1: III-V epitaxy layer structure used for the integrated III-V on amorphous silicon laser.....	106

Bibliography

- [1] “Cisco white paper, ‘The Zettabyte Era: Trends and Analysis,’ (Cisco, 2015), http://www.cisco.com/c/en/us/solutions/collateral/service-provider/visual-networking-index-vni/VNI_Hyperconnectivity_WP.html.”
- [2] D. Mahgerefteh *et al.*, “Techno-Economic Comparison of Silicon Photonics and Multimode VCSELs,” *J. Light. Technol.*, vol. 34, no. 2, pp. 233–242, Jan. 2016.
- [3] L. A. Coldren, S. W. Corzine, and M. Mashanovitch, *Diode lasers and photonic integrated circuits*, 2nd ed. Hoboken, N.J: Wiley, 2012.
- [4] D. Kuchta *et al.*, “A 56.1 Gb/s NRZ modulated 850nm VCSEL-based optical link,” in *Optical Fiber Communication Conference*, 2013, p. OW1B–5.
- [5] B. E. A. Saleh and M. C. Teich, *Fundamentals of Photonics*, 2nd ed. Hoboken, N.J: Wiley, 2007.
- [6] A. Mekis *et al.*, “A Grating-Coupler-Enabled CMOS Photonics Platform,” *IEEE J. Sel. Top. Quantum Electron.*, vol. 17, no. 3, pp. 597–608, May 2011.
- [7] D. H. Lee, S. J. Choo, U. Jung, K. W. Lee, K. W. Kim, and J. H. Park, “Low-loss silicon waveguides with sidewall roughness reduction using a SiO₂ hard mask and fluorine-based dry etching,” *J. Micromechanics Microengineering*, vol. 25, no. 1, p. 015003, Jan. 2015.
- [8] C. Sciancalepore *et al.*, “Low-crosstalk fabrication-insensitive echelle grating demultiplexers,” *IEEE Photonics Technol. Lett.*, vol. 27, no. 5, pp. 494–497, 2014.
- [9] L. Virot, “Développement de photodiodes à avalanche en Ge sur Si pour la détection faible signal et grande vitesse,” Université Paris Sud-Paris XI, 2014.
- [10] N. B. Feilchenfeld *et al.*, “An integrated silicon photonics technology for O-band datacom,” in *2015 IEEE International Electron Devices Meeting (IEDM)*, 2015, pp. 25–7.
- [11] F. Boeuf *et al.*, “Silicon Photonics R&D and Manufacturing on 300-mm Wafer Platform,” *J. Light. Technol.*, vol. 34, no. 2, pp. 286–295, Jan. 2016.
- [12] A. Novack *et al.*, “A 30 GHz silicon photonic platform,” 2013, p. 878107.
- [13] C. Baudot *et al.*, “Introducing photonic devices for 40Gbits/s wavelength division multiplexing transceivers on 300-mm SOI wafers using CMOS processes,” in *SPIE OPTO*, 2014, p. 89880O–89880O.
- [14] A. E.-J. Lim *et al.*, “Review of Silicon Photonics Foundry Efforts,” *IEEE J. Sel. Top. Quantum Electron.*, vol. 20, no. 4, pp. 405–416, Jul. 2014.
- [15] D. Liang and J. E. Bowers, “Recent progress in lasers on silicon,” *Nat. Photonics*, vol. 4, no. 8, pp. 511–517, Aug. 2010.
- [16] H. Rong *et al.*, “A continuous-wave Raman silicon laser,” *Nature*, vol. 433, no. 7027, pp. 725–725, Feb. 2005.
- [17] Z. Fang, Q. Y. Chen, and C. Z. Zhao, “A review of recent progress in lasers on silicon,” *Opt. Laser Technol.*, vol. 46, pp. 103–110, Mar. 2013.
- [18] J. Liu, L. C. Kimerling, and J. Michel, “Monolithic Ge-on-Si lasers for large-scale electronic–photonic integration,” *Semicond. Sci. Technol.*, vol. 27, no. 9, p. 094006, Sep. 2012.
- [19] J. Liu *et al.*, “Tensile-strained, n-type Ge as a gain medium for monolithic laser integration on Si,” *Opt. Express*, vol. 15, no. 18, pp. 11272–11277, 2007.
- [20] R. E. Camacho-Aguilera *et al.*, “An electrically pumped germanium laser,” *Opt. Express*, vol. 20, no. 10, pp. 11316–11320, 2012.
- [21] R. Koerner *et al.*, “Electrically pumped lasing from Ge Fabry-Perot resonators on Si,” *Opt. Express*, vol. 23, no. 11, p. 14815, Jun. 2015.
- [22] R. Geiger, T. Zabel, and H. Sigg, “Group IV Direct Band Gap Photonics: Methods, Challenges, and Opportunities,” *Front. Mater.*, vol. 2, Jul. 2015.
- [23] D. Nam *et al.*, “Strained germanium thin film membrane on silicon substrate for optoelectronics,” *Opt. Express*, vol. 19, no. 27, pp. 25866–25872, 2011.
- [24] G. Hiblot, “Modélisation compacte de transistors MOSFETs à canal III-V et films minces pour applications CMOS avancées,” Université Grenoble Alpes, 2015.
- [25] S. Chen *et al.*, “Electrically pumped continuous-wave III–V quantum dot lasers on silicon,” *Nat. Photonics*, vol. 10, no. 5, pp. 307–311, Mar. 2016.
- [26] Z. Wang *et al.*, “Room-temperature InP distributed feedback laser array directly grown on silicon,” *Nat. Photonics*, vol. 9, no. 12, pp. 837–842, Oct. 2015.
- [27] T. Pinguet *et al.*, “Advanced silicon photonic transceivers,” in *2015 IEEE 12th International Conference on Group IV Photonics (GFP)*, 2015, pp. 21–22.

-
- [28] Y. Urino *et al.*, “First demonstration of high density optical interconnects integrated with lasers, optical modulators, and photodetectors on single silicon substrate,” *Opt. Express*, vol. 19, no. 26, pp. B159–B165, 2011.
 - [29] B. Snyder, B. Corbett, and P. O’Brien, “Hybrid Integration of the Wavelength-Tunable Laser With a Silicon Photonic Integrated Circuit,” *J. Light. Technol.*, vol. 31, no. 24, pp. 3934–3942, Dec. 2013.
 - [30] T. Suzuki *et al.*, “Simple Method for Integrating DFB Laser and Silicon Photonics Platform,” in *2015 IEEE 12th International Conference on Group IV Photonics (GFP)*, 2015, pp. 27–28.
 - [31] Y. Urino *et al.*, “Demonstration of 12.5-Gbps optical interconnects integrated with lasers, optical splitters, optical modulators and photodetectors on a single silicon substrate,” *Opt. Express*, vol. 20, no. 26, pp. B256–B263, 2012.
 - [32] N. Hatori *et al.*, “A Hybrid Integrated Light Source on a Silicon Platform Using a Trident Spot-Size Converter,” *J. Light. Technol.*, vol. 32, no. 7, pp. 1329–1336, Apr. 2014.
 - [33] Y. Urino *et al.*, “First Demonstration of Athermal Silicon Optical Interposers With Quantum Dot Lasers Operating up to 125 °C,” *J. Light. Technol.*, vol. 33, no. 6, pp. 1223–1229, Mar. 2015.
 - [34] S. Romero-Garcia, B. Marzban, F. Merget, Bin Shen, and J. Witzens, “Edge Couplers With Relaxed Alignment Tolerance for Pick-and-Place Hybrid Integration of III-V Lasers With SOI Waveguides,” *IEEE J. Sel. Top. Quantum Electron.*, vol. 20, no. 4, pp. 369–379, Jul. 2014.
 - [35] S. Yang *et al.*, “A single adiabatic microring-based laser in 220 nm silicon-on-insulator,” *Opt. Express*, vol. 22, no. 1, p. 1172, Jan. 2014.
 - [36] A. J. Zilkie *et al.*, “Power-efficient III-V/Silicon external cavity DBR lasers,” *Opt. Express*, vol. 20, no. 21, pp. 23456–23462, 2012.
 - [37] D. Feng *et al.*, “High-Speed GeSi Electroabsorption Modulator on the SOI Waveguide Platform,” *IEEE J. Sel. Top. Quantum Electron.*, vol. 19, no. 6, pp. 64–73, Nov. 2013.
 - [38] S. Tanaka, S.-H. Jeong, S. Sekiguchi, T. Kurahashi, Y. Tanaka, and K. Morito, “High-output-power, single-wavelength silicon hybrid laser using precise flip-chip bonding technology,” *Opt. Express*, vol. 20, no. 27, pp. 28057–28069, 2012.
 - [39] S. Tanaka, T. Akiyama, S. Sekiguchi, and K. Morito, “Silicon Photonics Optical Transmitter Technology for Tb/s-class I/O Co-packaged with CPU,” *Fujitsu Sci Tech J*, vol. 50, no. 1, pp. 123–131, 2014.
 - [40] J. H. Lee *et al.*, “High power and widely tunable Si hybrid external-cavity laser for power efficient Si photonics WDM links,” *Opt. Express*, vol. 22, no. 7, p. 7678, Apr. 2014.
 - [41] J.-H. Lee *et al.*, “Demonstration of 12.2% wall plug efficiency in uncooled single mode external-cavity tunable Si/III-V hybrid laser,” *Opt. Express*, vol. 23, no. 9, p. 12079, May 2015.
 - [42] A. W. Fang, H. Park, O. Cohen, R. Jones, M. J. Paniccia, and J. E. Bowers, “Electrically pumped hybrid AlGaInAs-silicon evanescent laser,” *Opt. Express*, vol. 14, no. 20, pp. 9203–9210, 2006.
 - [43] A. W. Fang, E. Lively, Y.-H. Kuo, D. Liang, and J. E. Bowers, “A distributed feedback silicon evanescent laser,” *Opt. Express*, vol. 16, no. 7, pp. 4413–4419, 2008.
 - [44] H.-H. Chang *et al.*, “1310nm silicon evanescent laser,” *Opt. Express*, vol. 15, no. 18, pp. 11466–11471, 2007.
 - [45] X. Sun *et al.*, “Electrically pumped hybrid evanescent Si/InGaAsP lasers,” *Opt. Lett.*, vol. 34, no. 9, pp. 1345–1347, 2009.
 - [46] B. B. Bakir, A. Descos, N. Olivier, D. Bordel, P. Grosse, and J.-M. Fedeli, “Hybrid Silicon/III-V laser sources based on adiabatic mode transformers,” in *Photonics West*, 2012.
 - [47] A. W. Fang *et al.*, “Integrated AlGaInAs-silicon evanescent race track laser and photodetector,” *Opt. Express*, vol. 15, no. 5, pp. 2315–2322, 2007.
 - [48] A. W. Fang *et al.*, “A Distributed Bragg Reflector Silicon Evanescent Laser,” *IEEE Photonics Technol. Lett.*, vol. 20, no. 20, pp. 1667–1669, Oct. 2008.
 - [49] C. Zhang, S. Srinivasan, Y. Tang, M. J. R. Heck, M. L. Davenport, and J. E. Bowers, “Low threshold and high speed short cavity distributed feedback hybrid silicon lasers,” *Opt. Express*, vol. 22, no. 9, p. 10202, May 2014.
 - [50] H. Duprez *et al.*, “1310 nm hybrid InP/InGaAsP on silicon distributed feedback laser with high side-mode suppression ratio,” *Opt. Express*, vol. 23, no. 7, p. 8489, Apr. 2015.
 - [51] B. Ben Bakir *et al.*, “Heterogeneously Integrated III-V on Silicon Lasers,” *ECS Trans.*, vol. 64, no. 5, pp. 211–223, Aug. 2014.
 - [52] S. Keyvaninia *et al.*, “Demonstration of a heterogeneously integrated III-V/SOI single wavelength tunable laser,” *Opt. Express*, vol. 21, no. 3, pp. 3784–3792, 2013.
 - [53] H. Duprez, C. Jany, C. Seassal, and B. Ben Bakir, “Highly tunable heterogeneously integrated III-V on silicon sampled-grating distributed Bragg reflector lasers operating in the O-band,” *Opt. Express*, vol. 24, no. 18, p. 20895, Sep. 2016.

-
- [54] J. Van Campenhout *et al.*, “Electrically pumped InP-based microdisk lasers integrated with a nanophotonic silicon-on-insulator waveguide circuit,” *Opt. Express*, vol. 15, no. 11, pp. 6744–6749, 2007.
 - [55] D. Liang *et al.*, “Electrically-pumped compact hybrid silicon microring lasers for optical interconnects,” *Opt. Express*, vol. 17, no. 22, pp. 20355–20364, 2009.
 - [56] Y. Cao *et al.*, “Hybrid III–V/silicon laser with laterally coupled Bragg grating,” *Opt. Express*, vol. 23, no. 7, p. 88098817, 2015.
 - [57] A. Descos *et al.*, “Heterogeneously integrated III-V/Si distributed Bragg reflector laser with adiabatic coupling,” in *IET Conference Proceedings*, 2013.
 - [58] S. Srinivasan, A. W. Fang, D. Liang, J. Peters, B. Kaye, and J. E. Bowers, “Design of phase-shifted hybrid silicon distributed feedback lasers,” *Opt. Express*, vol. 19, no. 10, pp. 9255–9261, 2011.
 - [59] M. Lamponi *et al.*, “Low-Threshold Heterogeneously Integrated InP/SOI Lasers With a Double Adiabatic Taper Coupler,” *IEEE Photonics Technol. Lett.*, vol. 24, no. 1, pp. 76–78, Jan. 2012.
 - [60] S. Keyvaninia *et al.*, “Heterogeneously integrated III-V/silicon distributed feedback lasers,” *Opt. Lett.*, vol. 38, no. 24, p. 5434, Dec. 2013.
 - [61] S. Stankovic, R. Jones, M. N. Sysak, J. M. Heck, G. Roelkens, and D. Van Thourhout, “1310-nm Hybrid III-V/Si Fabry-Perot Laser Based on Adhesive Bonding,” *IEEE Photonics Technol. Lett.*, vol. 23, no. 23, pp. 1781–1783, Dec. 2011.
 - [62] M. Lamponi *et al.*, “Heterogeneously integrated InP/SOI laser using double tapered single-mode waveguides through adhesive die to wafer bonding,” in *2010 7th IEEE international conference on group IV photonics (GFP)*, 2010, pp. 22–24.
 - [63] L. Yuan *et al.*, “Hybrid InGaAsP-Si Evanescent Laser by Selective-Area Metal-Bonding Method,” *IEEE Photonics Technol. Lett.*, vol. 25, no. 12, pp. 1180–1183, Jun. 2013.
 - [64] H. Yu *et al.*, “1550-nm Evanescent Hybrid InGaAsP-Si Laser With Buried Ridge Stripe Structure,” *IEEE Photonics Technol. Lett.*, vol. 28, no. 10, pp. 1146–1149, May 2016.
 - [65] M. Li *et al.*, “A Hybrid Single-Mode Laser Based on Slotted Silicon Waveguides,” *IEEE Photonics Technol. Lett.*, pp. 1–1, 2016.
 - [66] T. Creazzo *et al.*, “Integrated tunable CMOS laser,” *Opt. Express*, vol. 21, no. 23, p. 28048, Nov. 2013.
 - [67] H. Chaouch *et al.*, “High Power, Narrow Linewidth, Low Noise, Integrated CMOS Tunable Laser for Long Haul Coherent Applications,” in *ECOC*, Cannes, 2014.
 - [68] E. Marchena *et al.*, “Integrated tunable CMOS laser for Si photonics,” in *Optical Fiber Communication Conference*, 2013, p. PDP5C–7.
 - [69] G. Roelkens *et al.*, “III-V-on-Silicon Photonic Devices for Optical Communication and Sensing,” *Photonics*, vol. 2, no. 3, pp. 969–1004, Sep. 2015.
 - [70] X. Luo *et al.*, “High-Throughput Multiple Dies-to-Wafer Bonding Technology and III/V-on-Si Hybrid Lasers for Heterogeneous Integration of Optoelectronic Integrated Circuits,” *Front. Mater.*, vol. 2, Apr. 2015.
 - [71] T. Komljenovic *et al.*, “Heterogeneous Silicon Photonic Integrated Circuits,” *J. Light. Technol.*, vol. 34, no. 1, pp. 20–35, Jan. 2016.
 - [72] J. Justice, C. Bower, M. Meitl, M. B. Mooney, M. A. Gubbins, and B. Corbett, “Wafer-scale integration of group III–V lasers on silicon using transfer printing of epitaxial layers,” *Nat. Photonics*, vol. 6, no. 9, pp. 612–616, Aug. 2012.
 - [73] H. Park *et al.*, “Device and Integration Technology for Silicon Photonic Transmitters,” *IEEE J. Sel. Top. Quantum Electron.*, vol. 17, no. 3, pp. 671–688, May 2011.
 - [74] L. Liu *et al.*, “Compact multiwavelength laser source based on cascaded InP-microdisks coupled to one SOI waveguide,” in *Optical Fiber Communication Conference*, 2008, p. OWQ3.
 - [75] J. Van Campenhout *et al.*, “Design and Optimization of Electrically Injected InP-Based Microdisk Lasers Integrated on and Coupled to a SOI Waveguide Circuit,” *J. Light. Technol.*, vol. 26, no. 1, pp. 52–63, Jan. 2008.
 - [76] G. Roelkens *et al.*, “III-V/silicon photonics for on-chip and intra-chip optical interconnects,” *Laser Photonics Rev.*, vol. 4, no. 6, pp. 751–779, Nov. 2010.
 - [77] J. Van Campenhout *et al.*, “Low-Footprint Optical Interconnect on an SOI Chip Through Heterogeneous Integration of InP-Based Microdisk Lasers and Microdetectors,” *IEEE Photonics Technol. Lett.*, vol. 21, no. 8, pp. 522–524, Apr. 2009.
 - [78] A. Yariv and X. Sun, “Supermode Si/III-V hybrid lasers, optical amplifiers and modulators: A proposal and analysis,” *Opt. Express*, vol. 15, no. 15, pp. 9147–9151, 2007.
 - [79] X. Sun, H.-C. Liu, and A. Yariv, “Adiabaticity criterion and the shortest adiabatic mode transformer in a coupled-waveguide system,” *Opt. Lett.*, vol. 34, no. 3, pp. 280–282, 2009.
 - [80] X. Sun, “Supermode Si/III–V lasers and circular Bragg lasers,” Citeseer, 2010.

-
- [81] S. Matsuo, T. Fujii, K. Hasebe, K. Takeda, T. Sato, and T. Kakitsuka, "Directly Modulated DFB Laser on SiO₂/Si Substrate for Datacenter Networks," *J. Light. Technol.*, vol. 33, no. 6, pp. 1217–1222, Mar. 2015.
 - [82] T. Fujii, T. Sato, K. Takeda, K. Hasebe, T. Kakitsuka, and S. Matsuo, "Epitaxial growth of InP to bury directly bonded thin active layer on SiO₂/Si substrate for fabricating distributed feedback lasers on silicon," *IET Optoelectron.*, vol. 9, no. 4, pp. 151–157, Aug. 2015.
 - [83] G. Kurczveil, P. Pintus, M. J. R. Heck, J. D. Peters, and J. E. Bowers, "Characterization of Insertion Loss and Back Reflection in Passive Hybrid Silicon Tapers," *IEEE Photonics J.*, vol. 5, no. 2, pp. 6600410–6600410, Apr. 2013.
 - [84] A. Alduino *et al.*, "Demonstration of a High Speed 4-Channel Integrated Silicon Photonics WDM Link with Hybrid Silicon Lasers," in *IPRPS*, 2010, p. PDIWI5.
 - [85] H.-F. Liu, "Demonstration of a 4X12.5 Gb/s fully integrated silicon photonic link," in *Microoptics conference*, Sendai, 2011, p. G-1.
 - [86] G.-H. Duan *et al.*, "10 Gb/s integrated tunable hybrid III-V/si laser and silicon mach-zehnder modulator," in *European Conference and Exhibition on Optical Communication*, 2012, p. Tu-4.
 - [87] A. Ramaswamy *et al.*, "A WDM 4x28Gbps Integrated Silicon Photonic Transmitter driven by 32nm CMOS driver ICs," in *Optical Fiber Communication Conference*, 2015, p. Th5B.5.
 - [88] H. Park *et al.*, "Heterogeneous integration of silicon photonic devices and integrated circuits," in *Conference on Lasers and Electro-Optics/Pacific Rim*, 2015, p. 25J3_2.
 - [89] A. Abbasi *et al.*, "28 Gb/s direct modulation heterogeneously integrated InP/Si DFB laser," in *Proc. Eur. Conf. Opt. Commun.(ECOC)*, 2015.
 - [90] G. Ghione, *Semiconductor Devices for High-Speed Optoelectronics*. Leiden: Cambridge University Press, 2009.
 - [91] G. T. Reed, G. Mashanovich, F. Y. Gardes, and D. J. Thomson, "Silicon optical modulators," *Nat. Photonics*, vol. 4, no. 8, pp. 518–526, Aug. 2010.
 - [92] L. Pavesi and G. Guillot, Eds., *Optical interconnects: the silicon approach*. Berlin ; New York: Springer, 2006.
 - [93] R. Soref and B. Bennett, "Electrooptical effects in silicon," *IEEE J. Quantum Electron.*, vol. 23, no. 1, pp. 123–129, 1987.
 - [94] B. J. Frey, D. B. Leviton, and T. J. Madison, "Temperature-dependent refractive index of silicon and germanium," in *SPIE Astronomical Telescopes+ Instrumentation*, 2006, p. 62732J–62732J.
 - [95] U. Fisher, T. Zinke, B. Schuppert, and K. Petermann, "Singlemode optical switches based on SOI waveguides with large cross-section," *Electron. Lett.*, vol. 30, no. 5, pp. 406–408, 1994.
 - [96] Y. Tang, J. D. Peters, and J. E. Bowers, "Over 67 GHz bandwidth hybrid silicon electroabsorption modulator with asymmetric segmented electrode for 1.3 μm transmission," *Opt. Express*, vol. 20, no. 10, pp. 11529–11535, 2012.
 - [97] Y. Tang, H.-W. Chen, S. Jain, J. D. Peters, U. Westergren, and J. E. Bowers, "50 Gb/s hybrid silicon traveling-wave electroabsorption modulator," *Opt. Express*, vol. 19, no. 7, pp. 5811–5816, 2011.
 - [98] Y. Kuo, H.-W. Chen, and J. E. Bowers, "High speed hybrid silicon evanescent electroabsorption modulator," *Opt. Express*, vol. 16, no. 13, pp. 9936–9941, 2008.
 - [99] Yi-Jen Chiu, Tsu-Hsiu Wu, Wen-Chin Cheng, F. J. Lin, and J. E. Bowers, "Enhanced performance in traveling-wave electroabsorption Modulators based on undercut-etching the active-region," *IEEE Photonics Technol. Lett.*, vol. 17, no. 10, pp. 2065–2067, Oct. 2005.
 - [100] M. N. Sysak, J. O. Anthes, J. E. Bowers, O. Raday, and R. Jones, "Integration of hybrid silicon lasers and electroabsorption modulators," *Opt. Express*, vol. 16, no. 17, pp. 12478–12486, 2008.
 - [101] S. R. Jain, M. N. Sysak, G. Kurczveil, and J. E. Bowers, "Integrated hybrid silicon DFB laser-EAM array using quantum well intermixing," *Opt. Express*, vol. 19, no. 14, pp. 13692–13699, 2011.
 - [102] S. Jain, M. N. Sysak, M. Piels, and J. E. Bowers, "Hybrid silicon transmitter using quantum well intermixing," in *Optical Fiber Communication Conference*, 2013, p. OTh1D-2.
 - [103] S. A. Srinivasan *et al.*, "56 Gb/s Germanium Waveguide Electro-Absorption Modulator," *J. Light. Technol.*, vol. 34, no. 2, pp. 419–424, Jan. 2016.
 - [104] L. Lever *et al.*, "Modulation of the absorption coefficient at 1.3 μm in Ge/SiGe multiple quantum well heterostructures on silicon," *Opt. Lett.*, vol. 36, no. 21, pp. 4158–4160, 2011.
 - [105] L. Vivien and L. Pavesi, *Handbook of Silicon Photonics*. CRC Press, 2013.
 - [106] M. Nedeljkovic, R. Soref, and G. Z. Mashanovich, "Free-Carrier Electrorefraction and Electroabsorption Modulation Predictions for Silicon Over the 1-14- μm Infrared Wavelength Range," *IEEE Photonics J.*, vol. 3, no. 6, pp. 1171–1180, Dec. 2011.
 - [107] O. Dubray *et al.*, "Electro-optical ring modulator: an ultra-compact model for the comparison and optimization of PN, PIN, and capacitive junction," *IEEE J. Sel. Top. Quantum Electron.*, pp. 1–1, 2016.

-
- [108] T. Baba *et al.*, “50-Gb/s ring-resonator-based silicon modulator,” *Opt. Express*, vol. 21, no. 10, p. 11869, May 2013.
 - [109] S. Akiyama *et al.*, “Compact PIN-Diode-Based Silicon Modulator Using Side-Wall-Grating Waveguide,” *IEEE J. Sel. Top. Quantum Electron.*, vol. 19, no. 6, pp. 74–84, Nov. 2013.
 - [110] S. Kupijai *et al.*, “25 Gb/s Silicon Photonics Interconnect Using a Transmitter Based on a Node-Matched-Diode Modulator,” *J. Light. Technol.*, vol. 34, no. 12, pp. 2920–2923, Jun. 2016.
 - [111] T. Baba, S. Akiyama, M. Imai, and T. Usuki, “25-Gb/s broadband silicon modulator with 0.31-V·cm $V\pi$ L based on forward-biased PIN diodes embedded with passive equalizer,” *Opt. Express*, vol. 23, no. 26, p. 32950, Dec. 2015.
 - [112] M. Webster *et al.*, “An efficient MOS-capacitor based silicon modulator and CMOS drivers for optical transmitters,” in *11th International Conference on Group IV Photonics (GFP)*, Paris, 2014, pp. 1–2.
 - [113] A. Abraham, S. Olivier, D. Marris-Morini, and L. Vivien, “Evaluation of the performances of a silicon optical modulator based on a silicon-oxide-silicon capacitor,” in *11th International Conference on Group IV Photonics (GFP)*, 2014, pp. 3–4.
 - [114] L. Liao *et al.*, “High speed silicon Mach-Zehnder modulator,” *Opt. Express*, vol. 13, no. 8, pp. 3129–3135, 2005.
 - [115] M. Sodagar *et al.*, “Compact, 15 Gb/s electro-optic modulator through carrier accumulation in a hybrid Si/SiO₂/Si microdisk,” *Opt. Express*, vol. 23, no. 22, p. 28306, Nov. 2015.
 - [116] X. Tu *et al.*, “Silicon optical modulator with shield coplanar waveguide electrodes,” *Opt. Express*, vol. 22, no. 19, p. 23724, Sep. 2014.
 - [117] X. Tu *et al.*, “50-Gb/s silicon Mach-Zehnder interferometer-based optical modulator with only 1.3 V pp driving voltages,” in *Electronics Packaging Technology Conference (EPTC), 2014 IEEE 16th*, 2014, pp. 851–854.
 - [118] X. Tu *et al.*, “50-Gb/s silicon optical modulator with traveling-wave electrodes,” *Opt. Express*, vol. 21, no. 10, p. 12776, May 2013.
 - [119] M. Ziebell *et al.*, “40 Gbit/s low-loss silicon optical modulator based on a pipin diode,” *Opt. Express*, vol. 20, no. 10, pp. 10591–10596, 2012.
 - [120] Hao Xu, Xianyao Li, Xi Xiao, Zhiyong Li, Yude Yu, and Jinzhong Yu, “Demonstration and Characterization of High-Speed Silicon Depletion-Mode Mach-Zehnder Modulators,” *IEEE J. Sel. Top. Quantum Electron.*, vol. 20, no. 4, pp. 23–32, Jul. 2014.
 - [121] D. Marris-Morini *et al.*, “Low loss 40 Gbit/s silicon modulator based on interleaved junctions and fabricated on 300 mm SOI wafers,” *Opt. Express*, vol. 21, no. 19, p. 22471, Sep. 2013.
 - [122] H. Yu *et al.*, “Performance tradeoff between lateral and interdigitated doping patterns for high speed carrier-depletion based silicon modulators,” *Opt. Express*, vol. 20, no. 12, pp. 12926–12938, 2012.
 - [123] X. Xiao *et al.*, “High-speed, low-loss silicon Mach-Zehnder modulators with doping optimization,” *Opt. Express*, vol. 21, no. 4, pp. 4116–4125, 2013.
 - [124] J. Wang *et al.*, “Optimization and Demonstration of a Large-bandwidth Carrier-depletion Silicon Optical Modulator,” *J. Light. Technol.*, vol. 31, no. 24, pp. 4119–4125, Dec. 2013.
 - [125] P. Dong, L. Chen, and Y.-K. Chen, “High-speed Low-voltage Single-drive Push-pull Silicon Mach-Zehnder Modulators,” *Opt. Express*, vol. 20, no. 6, pp. 6163–6169, 2012.
 - [126] D. Patel *et al.*, “Design, analysis, and transmission system performance of a 41 GHz silicon photonic modulator,” *Opt. Express*, vol. 23, no. 11, p. 14263, Jun. 2015.
 - [127] R. Ding *et al.*, “High-Speed Silicon Modulator With Slow-Wave Electrodes and Fully Independent Differential Drive,” *J. Light. Technol.*, vol. 32, no. 12, pp. 2240–2247, Jun. 2014.
 - [128] M. Streshinsky *et al.*, “Low power 50 Gb/s silicon traveling wave Mach-Zehnder modulator near 1300 nm,” *Opt. Express*, vol. 21, no. 25, p. 30350, Dec. 2013.
 - [129] M. Ziebell, “Silicon Optical Transceiver for Local Access Networks,” Université Paris Sud-Paris XI, 2013.
 - [130] H. Yu and W. Bogaerts, “An Equivalent Circuit Model of the Traveling Wave Electrode for Carrier-Depletion-Based Silicon Optical Modulators,” *J. Light. Technol.*, vol. 30, no. 11, pp. 1602–1609, Jun. 2012.
 - [131] M. Pantouvaki, P. Verheyen, J. De Coster, G. Lepage, P. Absil, and J. Van Campenhout, “56Gb/s ring modulator on a 300mm silicon photonics platform,” in *Optical Communication (ECOC), 2015 European Conference on*, 2015, pp. 1–3.
 - [132] R. Dube-Demers *et al.*, “Analytical Modeling of Silicon Microring and Microdisk Modulators With Electrical and Optical Dynamics,” *J. Light. Technol.*, vol. 33, no. 20, pp. 4240–4252, Oct. 2015.
 - [133] P. Le Maître *et al.*, “Impact of process variability of active ring resonators in a 300mm silicon photonic platform,” in *Optical Communication (ECOC), 2015 European Conference on*, 2015, pp. 1–3.
 - [134] H. Xu, X.-Y. Li, X. Xiao, Z.-Y. Li, Y.-D. Yu, and J.-Z. Yu, “High-speed and broad optical bandwidth silicon modulator,” *Chin. Phys. B*, vol. 22, no. 11, p. 114212, Nov. 2013.

-
- [135] Z. Zhou, B. Yin, Q. Deng, X. Li, and J. Cui, "Lowering the energy consumption in silicon photonic devices and systems [Invited]," *Photonics Res.*, vol. 3, no. 5, p. B28, Oct. 2015.
 - [136] Y. Liu, S. Dunham, T. Baehr-Jones, A. E.-J. Lim, G.-Q. Lo, and M. Hochberg, "Ultra-Responsive Phase Shifters for Depletion Mode Silicon Modulators," *J. Light. Technol.*, vol. 31, no. 23, pp. 3787–3793, Dec. 2013.
 - [137] G. Rasigade, "Modulateur optique haute-fréquence sur substrat silicium-sur-isolant," Université Paris Sud-Paris XI, 2010.
 - [138] W. R. Eisenstadt and E. Yungseon, "S-Parameter-Based IC Interconnect Transmission Line Characterization," *IEEE Trans. Compon. Hybrids Manuf. Technol.*, vol. 15, no. 4, pp. 483–490, 1992.
 - [139] H. Jayatilaka, W. D. Sacher, and J. K. S. Poon, "Analytical Model and Fringing-Field Parasitics of Carrier-Depletion Silicon-on-Insulator Optical Modulation Diodes," *IEEE Photonics J.*, vol. 5, no. 1, pp. 2200211–2200211, Feb. 2013.
 - [140] S. M. Sze and K. K. Ng, *Physics of semiconductor devices*, 3rd ed. Hoboken, N.J: Wiley-Interscience, 2007.
 - [141] C. R. Neve *et al.*, "RF and linear performance of commercial 200 mm trap-rich HR-SOI wafers for SoC applications," in *Silicon Monolithic Integrated Circuits in RF Systems (SiRF)*, 2013 IEEE 13th Topical Meeting on, 2013, pp. 15–17.
 - [142] S. Sharif Azadeh, J. Müller, F. Merget, S. Romero-García, B. Shen, and J. Witzens, "Advances in silicon photonics segmented electrode Mach-Zehnder modulators and peaking enhanced resonant devices," 2014, p. 928817.
 - [143] F. Merget *et al.*, "Silicon photonics plasma-modulators with advanced transmission line design," *Opt. Express*, vol. 21, no. 17, p. 19593, Aug. 2013.
 - [144] P. Dong *et al.*, "Novel integration technique for silicon/III-V hybrid laser," *Opt. Express*, vol. 22, no. 22, p. 26854, Nov. 2014.
 - [145] A. Katz *et al.*, "Pt/Ti/n-InP nonalloyed ohmic contacts formed by rapid thermal processing," *J. Appl. Phys.*, vol. 67, no. 8, p. 3872, 1990.
 - [146] S. Jain, M. Sysak, M. Swaidan, and J. Bowers, "Silicon fab-compatible contacts to n-InP and p-InGaAs for photonic applications," *Appl. Phys. Lett.*, vol. 100, no. 20, p. 201103, 2012.
 - [147] A. Katz *et al.*, "Pt/Ti/p-In_{0.53}Ga_{0.47}As low-resistance nonalloyed ohmic contact formed by rapid thermal processing," *Appl. Phys. Lett.*, vol. 54, no. 23, p. 2306, 1989.
 - [148] E. F. Chor, W. K. Chong, and C. H. Heng, "Alternative (Pd,Ti,Au) contacts to (Pt,Ti,Au) contacts for In_{0.53}Ga_{0.47}As," *J. Appl. Phys.*, vol. 84, no. 5, p. 2977, 1998.
 - [149] B. B. Bakir *et al.*, "Electrically driven hybrid Si/III-V Fabry-Pérot lasers based on adiabatic mode transformers," *Opt. Express*, vol. 19, no. 11, pp. 10317–10325, 2011.
 - [150] A. Katz, B. E. Weir, and W. C. Dautremont-Smith, "Au/Pt/Ti contacts to p-In_{0.53}Ga_{0.47}As and n-InP layers formed by a single metallization common step and rapid thermal processing," *J. Appl. Phys.*, vol. 68, no. 3, p. 1123, 1990.
 - [151] A. Descos, "Conception, fabrication et réalisation de sources lasers hybrides III-V sur silicium," Ecully, Ecole centrale de Lyon, 2014.
 - [152] D. J. Thomson *et al.*, "High Performance Mach-Zehnder-Based Silicon Optical Modulators," *IEEE J. Sel. Top. Quantum Electron.*, vol. 19, no. 6, pp. 85–94, Nov. 2013.
 - [153] J. Pu *et al.*, "Heterogeneous Integrated III–V Laser on Thin SOI With Single-Stage Adiabatic Coupler: Device Realization and Performance Analysis," *IEEE J. Sel. Top. Quantum Electron.*, vol. 21, no. 6, pp. 369–376, Nov. 2015.
 - [154] Y. De Koninck, G. Roelkens, and R. Baets, "Electrically pumped 1550 nm single mode III-V-on-silicon laser with resonant grating cavity mirrors: Electrically pumped 1550 nm single mode," *Laser Photonics Rev.*, vol. 9, no. 2, pp. L6–L10, Mar. 2015.
 - [155] Q. Wang *et al.*, "Heterogeneous Si/III-V integration and the optical vertical interconnect access," *Opt. Express*, vol. 20, no. 15, pp. 16745–16756, 2012.
 - [156] Q. Huang, J. Cheng, L. Liu, Y. Tang, and S. He, "Ultracompact adiabatic tapered coupler for the Si/III-V heterogeneous integration," in *PIERS Proceedings*, 2014, pp. 920–924.
 - [157] S. Zhu, Q. Fang, M. B. Yu, G. Q. Lo, and D. L. Kwong, "Propagation losses in undoped and n-doped polycrystalline silicon wire waveguides," *Opt. Express*, vol. 17, no. 23, pp. 20891–20899, 2009.
 - [158] S. Zhu, G. Q. Lo, J. D. Ye, and D. L. Kwong, "Influence of RTA and LTA on the Optical Propagation Loss in Polycrystalline Silicon Wire Waveguides," *IEEE Photonics Technol. Lett.*, vol. 22, no. 7, pp. 480–482, Apr. 2010.
 - [159] J. Poortmans and V. Arkhipov, Eds., *Thin film solar cells: fabrication, characterization and applications*. Chichester, England ; Hoboken, NJ: Wiley, 2007.
 - [160] F. G. Della Corte and S. Rao, "Use of Amorphous Silicon for Active Photonic Devices," *IEEE Trans. Electron Devices*, vol. 60, no. 5, pp. 1495–1505, May 2013.

-
- [161] A. Harke, M. Krause, and J. Mueller, "Low-loss singlemode amorphous silicon waveguides," *Electron. Lett.*, vol. 41, no. 25, 2005.
 - [162] S. Zhu, G. Q. Lo, W. Li, and D. L. Kwong, "Effect of cladding layer and subsequent heat treatment on hydrogenated amorphous silicon waveguides," *Opt. Express*, vol. 20, no. 21, pp. 23676–23683, 2012.
 - [163] S. Zhu, G. Q. Lo, and D. L. Kwong, "Low-loss amorphous silicon wire waveguide for integrated photonics: effect of fabrication process and the thermal stability," *Opt. Express*, vol. 18, no. 24, pp. 25283–25291, 2010.
 - [164] S. K. Selvaraja *et al.*, "Low-loss amorphous silicon-on-insulator technology for photonic integrated circuitry," *Opt. Commun.*, vol. 282, no. 9, pp. 1767–1770, 2009.
 - [165] T. Lipka, L. Moldenhauer, J. Müller, and H. K. Trieu, "Photonic integrated circuit components based on amorphous silicon-on-insulator technology," *Photonics Res.*, vol. 4, no. 3, p. 126, Jun. 2016.
 - [166] C. Grillet *et al.*, "Amorphous silicon nanowires combining high nonlinearity, FOM and optical stability," *Opt. Express*, vol. 20, no. 20, pp. 22609–22615, 2012.
 - [167] H. Ghafouri-Shiraz, *Distributed feedback laser diodes and optical tunable filters*. Chichester: Wiley, 2003.
 - [168] D. Vermeulen *et al.*, "High-efficiency fiber-to-chip grating couplers realized using an advanced CMOS-compatible silicon-on-insulator platform," *Opt. Express*, vol. 18, no. 17, pp. 18278–18283, 2010.
 - [169] "<https://static1.squarespace.com/static/55ae48f4e4b0d98862c1d3c7/t/578e71229de4bb5b172a1fd3/1468952923570/2016-06-09+AIM+meeting+-+LETI+-+ChristopheKOPP.pdf>."
 - [170] J. Durel *et al.*, "Realization of back-side heterogeneous hybrid III-V/Si DBR lasers for silicon photonics," in *Photonics West*, San Francisco, 2016.
 - [171] W. S. Zaoui, A. Kunze, W. Vogel, M. Berroth, J. Butschke, and F. Letzkus, "CMOS-compatible nonuniform grating coupler with 86% coupling efficiency," in *IET Conference Proceedings*, 2013.
 - [172] C. Baudot *et al.*, "Low cost 300mm double-SOI substrate for low insertion loss 1D & 2D grating couplers," in *11th International Conference on Group IV Photonics (GFP)*, 2014.
 - [173] Q. I. Nan *et al.*, "A 25Gb/s, 520mW, 6.4 Vpp silicon-photonic Mach-Zehnder modulator with distributed driver in CMOS," in *Optical Fiber Communication Conference*, 2015, p. W1B–3.
 - [174] G. Denoyer *et al.*, "Hybrid Silicon Photonic Circuits and Transceiver for 50 Gb/s NRZ Transmission Over Single-Mode Fiber," *J. Light. Technol.*, vol. 33, no. 6, pp. 1247–1254, Mar. 2015.
 - [175] E. Temporiti, G. Minoia, M. Repossi, D. Baldi, A. Ghilioni, and F. Svelto, "23.4 A 56Gb/s 300mW silicon-photonics transmitter in 3D-integrated PIC25G and 55nm BiCMOS technologies," in *2016 IEEE International Solid-State Circuits Conference (ISSCC)*, 2016, pp. 404–405.
 - [176] M. Cignoli *et al.*, "22.9 A 1310nm 3D-integrated silicon photonics Mach-Zehnder-based transmitter with 275mW multistage CMOS driver achieving 6dB extinction ratio at 25Gb/s," in *2015 IEEE International Solid-State Circuits Conference-(ISSCC) Digest of Technical Papers*, 2015, pp. 1–3.
 - [177] J. Prades, A. Ghiotto, D. Pache, and E. Kerhervé, "High speed phase modulator driver unit in 55 nm SiGe BiCMOS for a single-channel 100 Gb/s NRZ silicon photonic modulator," in *Microwave Integrated Circuits Conference (EuMIC), 2015 10th European*, 2015, pp. 341–344.
 - [178] Chong Zhang, Di Liang, G. Kurczveil, J. E. Bowers, and R. G. Beausoleil, "Thermal Management of Hybrid Silicon Ring Lasers for High Temperature Operation," *IEEE J. Sel. Top. Quantum Electron.*, vol. 21, no. 6, pp. 385–391, Nov. 2015.
 - [179] M. N. Sysak, H. Park, A. W. Fang, O. Raday, J. E. Bowers, and R. Jones, "Reduction of hybrid silicon laser thermal impedance using Poly Si thermal shunts," in *Proc. Conf. Opt. Fiber Commun., Collocated Nat. Fiber Opt. Eng. Conf.*, 2011, pp. 1–3.
 - [180] A. Y. Liu *et al.*, "High performance continuous wave 1.3 μm quantum dot lasers on silicon," *Appl. Phys. Lett.*, vol. 104, no. 4, p. 041104, Jan. 2014.
 - [181] G. Kurczveil, D. Liang, M. Fiorentino, and R. G. Beausoleil, "Robust hybrid quantum dot laser for integrated silicon photonics," *Opt. Express*, vol. 24, no. 14, p. 16167, Jul. 2016.
 - [182] "<http://www.ieee802.org/3/bs/index.html>."
 - [183] D. Patel, A. Samani, V. Veerasubramanian, S. Ghosh, and D. V. Plant, "Silicon Photonic Segmented Modulator-Based Electro-Optic DAC for 100 Gb/s PAM-4 Generation," *IEEE Photonics Technol. Lett.*, vol. 27, no. 23, pp. 2433–2436, Dec. 2015.
 - [184] R. Soref, "The Past, Present, and Future of Silicon Photonics," *IEEE J. Sel. Top. Quantum Electron.*, vol. 12, no. 6, pp. 1678–1687, Nov. 2006.
 - [185] D. Thomson *et al.*, "Roadmap on silicon photonics," *J. Opt.*, vol. 18, no. 7, p. 073003, Jul. 2016.
 - [186] G. Morthier, T. Spuesens, P. Mechet, G. Roelkens, and D. Van Thourhout, "InP Microdisk Lasers Integrated on Si for Optical Interconnects," *IEEE J. Sel. Top. Quantum Electron.*, vol. 21, no. 6, pp. 359–368, Nov. 2015.

Author list of publications

Journals:

First author:

T. Ferrotti, H. Duprez, C. Jany, A. Chantre, C. Seassal and B. Ben Bakir, "O-Band III-V-on-Amorphous-Silicon Lasers Integrated With a Surface Grating Coupler," IEEE Photonics Technology Letters, 28(18), p. 1944-1947 (2016).

T. Ferrotti, B. Blampey, C. Jany, H. Duprez, A. Chantre, F. Boeuf, C. Seassal and B. Ben Bakir, "Co-integrated 1.3 μ m Hybrid III-V/silicon Tunable Laser and Silicon Mach-Zehnder Modulator Operating at 25Gb/s," Optics Express, 24(26), pp. 30379-30401 (2016).

Co-author:

H. Duprez, A. Descos, **T. Ferrotti**, C. Sciancalepore, C. Jany, K. Hassan, C. Seassal, S. Menezo, and B. Ben Bakir, "1310 nm hybrid InP/InGaAsP on silicon distributed feedback laser with high side-mode suppression ratio," Optics Express, 23(7), p.8489-8497 (2015).

K. Hassan, C. Sciancalepore, J. Harduin, **T. Ferrotti**, S. Menezo, and B. Ben Bakir, "Toward athermal silicon-on-insulator (de)multiplexers in the O-band," Optics Express, 40(11), p.2641-2644 (2015).

Conference proceedings:

First author:

T. Ferrotti, A. Chantre, B. Blampey, H. Duprez, F. Milesi, A. Myko, C. Sciancalepore, K. Hassan, J. Harduin, C. Baudot, S. Menezo, F. Bœuf and B. Ben Bakir, "Power-efficient carrier-depletion SOI Mach-Zehnder modulators for 4x25 Gbit/s operation in the O-band," Photonics West, Proc. SPIE 9367, Silicon Photonics X, 93670D, San Francisco (2015).

T. Ferrotti, B. Blampey, H. Duprez, C. Jany, A. Chantre, F. Bœuf, C. Seassal and B. Ben Bakir, "1.3 μ m Hybrid III-V on Silicon Transmitter Operating at 25Gb/s," International Conference on Solid State Devices and Materials, paper C-3-02, Tsukuba (2016).

Co-author:

C. Sciancalepore, K. Hassan, **T. Ferrotti**, J. Harduin, H. Duprez, S. Menezo and B. Ben Bakir, "Low-loss adiabatically-tapered high-contrast gratings for slow-wave modulators on SOI," Photonics West, Proc. SPIE 9372, High Contrast Metastructures IV, 93720G, San Francisco (2015).

H. Duprez, A. Descos, **T. Ferrotti**, C. Jany, J. Harduin, A. Myko, C. Sciancalepore, C. Seassal and B. Ben Bakir, "Heterointegrated III-V/Si distributed feedback lasers," Photonics West, Proc. SPIE 9365, Integrated Optics: Devices, Materials, and Technologies XIX, 936506, San Francisco (2015).

C. Baudot, B. Szelag, N. Allouti, C. Comboroure, S. Bérard-Bergery, C. Vizioz, S. Barnola, F. Gays, D. Mariolle, **T. Ferrotti**, A. Souhaité, S. Brisson, C. Kopp and S. Menezo, "Progresses in 300mm DUV photolithography for the development of advanced silicon photonic devices," Advanced Lithography, Proc. SPIE 9426, Optical Microlithography XXVIII, 94260D, San Jose (2015).

H. Duprez, A. Descos, **T. Ferrotti**, J. Harduin, C. Jany, T. Card, A. Myko, C. Sciancalepore, S. Menezo, and B. Ben Bakir, "Heterogeneously Integrated III-V on Silicon Distributed Feedback Lasers at 1310 nm," Optical Fiber Communication Conference, paper Tu3I.6, Los Angeles (2015).

K. Hassan, C. Sciancalepore, J. Harduin, **T. Ferrotti**, S. Pauliac, C. Kopp, S. Menezo, B. Ben Bakir, R. John Lycett, D. F. G. Gallagher, J. A. Dallery and U. Weidenmueller, “Advances toward temperature-independent and fabrication-insensitive (de)multiplexer in the O-band,” IEEE 12th International Conference on Group IV Photonics (GFP), Vancouver (2015).

P. Le Maître, J. F. Carpentier, C. Baudot, N. Vuillet, A. Souhaité, J. B. Quélène, **T. Ferrotti** and F. Boeuf, “Impact of process variability of active ring resonators in a 300mm silicon photonic platform,” European Conference on Optical Communication (ECOC), Valencia (2015).

J. Durel, **T. Ferrotti**, A. Chantre, S. Cremer, J. Harduin, S. Bernabé, C. Kopp, F. Boeuf, B. Ben Bakir and J. E. Broquin, “Realization of back-side heterogeneous hybrid III-V/Si DBR lasers for silicon photonics,” Photonics West, Proc. SPIE 9750, Integrated Optics: Devices, Materials, and Technologies XX, 97500O, San Francisco (2016).

B. Szelag, B. Blampey, **T. Ferrotti**, V. Reboud, K. Hassan, S. Malhouitre, G. Grand, D. Fowler, S. Brision, T. Bria, G. Rabillé, P. Briancaeu, J. M. Hartmann, V. Hugues, A. Myko, F. Elleboode, F. Gays, J. M. Fédéli and C. Kopp, “Multiple wavelength silicon photonic 200 mm R+D platform for 25Gb/s and above applications,” Photonics Europe, Proc. SPIE 9891, Silicon Photonics and Photonic Integrated Circuits V, 98911C, Brussels (2016).

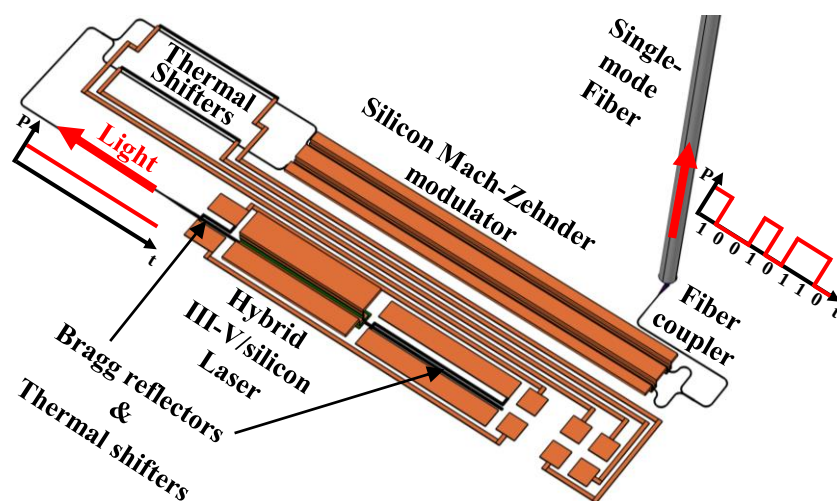
Patents:

T. Ferrotti, B. Ben Bakir, A. Chantre, S. Cremer and H. Duprez, “Laser device and process for fabricating such a laser device,” publication number US 20160056612 A1 (2016).

T. Ferrotti, B. Ben Bakir, A. Chantre, S. Cremer and H. Duprez, “Hybrid III-V Si laser architecture”, submission in progress.

Appendix A: French summary

Cette annexe contient un résumé en français du manuscrit, incluant les principaux résultats obtenus pendant la thèse. Pour commencer, le contexte ainsi que les objectifs de ce travail de thèse sont présentés, suivi des développements réalisés sur les modulateurs silicium de type Mach-Zehnder à déplétion de porteurs. Ensuite, le cœur de ce travail de thèse, portant sur la co-intégration du modulateur et d'une source laser hybride III-V sur silicium est détaillé. Une solution basée sur le dépôt d'une couche de silicium amorphe afin d'améliorer l'intégration de la source laser dans une route de fabrication standard pour la photonique sur silicium est aussi proposée. Finalement, les conclusions pouvant être tirées de ce travail de thèse sont présentées, de même que les pistes possibles pour le poursuivre.



A-1. Introduction.....	138
A-2. Modulateur Silicium	144
A-3. Transmetteur hybride III-V sur silicium	155
A-4. Laser hybride III-V sur silicium amorphe.....	164
A-5. Conclusions et perspectives	169

A-1. Introduction

Depuis plusieurs années, la quantité de données échangées à travers le monde ne cesse d'augmenter, et il est peu probable que cette tendance s'arrête. Selon des prévisions de l'entreprise CISCO [1], le volume d'informations transmis chaque année par Internet est au-delà du zettabyte (10^{21}) et devrait encore doubler à l'horizon 2019. Pour gérer une telle quantité de données, des débits de transmission de données de plus en plus élevés sur des distances de plusieurs centaines de mètres (voire plusieurs kilomètres) sont nécessaires, surtout dans les centres de données. Bien que les câbles à base de cuivre soient encore utilisés pour les communications à haut débits pour des distances inférieures à trois mètres, leur forte atténuation ($\approx 1\text{dB/m}$) ne les rend pas aptes pour des distances plus longues. Ainsi, ils tendent à être remplacés pour des distances de plus en plus courtes par les systèmes photoniques, qui bénéficient d'une atténuation réduite dans les fibres optiques ($\approx 0.2\text{-}0.3\text{dB/km}$). De plus, le multiplexage en longueur d'onde permet d'augmenter considérablement les débits de transmission à l'aide d'une seule fibre. Cependant, le coût élevé de ces systèmes, dû aux matériaux utilisés ainsi qu'à leur mise en boîtier, limite leur utilisation si celle-ci n'est pas indispensable.

Dans ce contexte, la photonique sur silicium est vue comme la solution pour obtenir des systèmes optiques de transmission de données performants et à coût réduit. En effet, le silicium étant utilisé depuis plusieurs décennies pour les composants microélectroniques, son utilisation donne accès à des capacités de fabrication à grande échelle et à rendement élevé. De plus, le silicium est transparent pour des longueurs d'onde supérieures à $1.2\mu\text{m}$, telles que celles utilisées pour les applications Datacom ($1.3\mu\text{m}$) et Telecom ($1.55\mu\text{m}$), et son contraste d'indice élevé avec son oxyde natif permet de concevoir des systèmes optiques efficaces et compacts. Pour finir, les électronique de commande ainsi que les amplificateurs à transimpédance peuvent être co-intégrés directement ou collés au-dessus des composants photoniques sur silicium, ce qui réduit la complexité d'intégration de ces systèmes.

Grâce à un intérêt croissant des industriels et de la communauté scientifique, la photonique sur silicium a connu un développement rapide au cours de la dernière décennie. Celui-ci a permis d'obtenir une vaste librairie de composants passifs et actifs intégrés sur silicium, nécessaires pour former des liens de transmission optique, tels que celui illustré sur la **figure A-1**. Ainsi, les fonctions de couplages entre circuits photoniques et fibre optique, de (dé)multiplexage optique, de modulation électro-optique à haut débit et de photo détection (à base de germanium) ont gagné en maturité, et sont produits à un niveau industriel [10], [11].

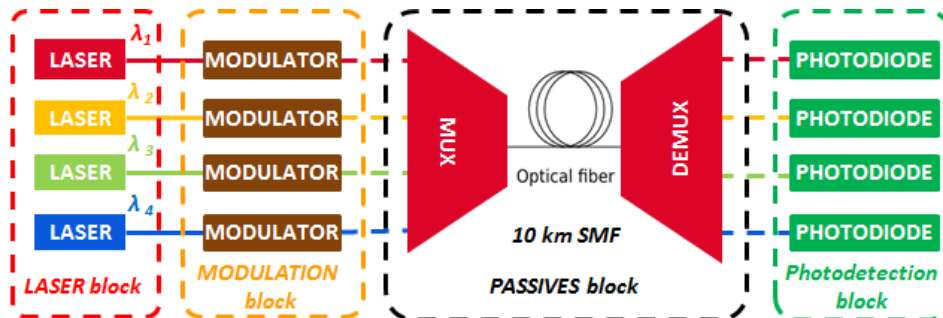


Figure A-1: Exemple de lien optique pour la norme de transmission 100GBASE-LR4.

Cependant, parmi les composants nécessaires pour former des liens de transmissions optiques, l'un d'entre eux demeure en retard en termes d'intégration sur silicium: la source laser, nécessaire pour générer les signaux optiques. En effet, le silicium étant un semi-conducteur à gap indirect, celui-ci est un mauvais candidat en termes d'émission de photons (émission spontanée) et d'amplification optique (émission stimulée). Or, ces deux mécanismes sont nécessaires au bon fonctionnement d'une source laser. En effet, ils requièrent à la fois la conservation de l'énergie et du moment durant la recombinaison d'un électron de la bande de conduction vers un trou de la bande de valence. Puisque le photon généré durant la recombinaison peut posséder suffisamment d'énergie, mais a un moment négligeable, la différence de moment entre le trou et l'électron doit elle aussi être négligeable. La **figure A-2** illustre cette condition pour des semi-conducteurs à gap direct et à gap indirect, qui est remplie assez aisément dans le premier cas, mais pas dans le second, où un phonon sera nécessaire pour conserver le moment. Ce type de recombinaison, faisant participer à la fois des électrons, des photons et des

phonons, ayant une probabilité extrêmement faible de se produire, les semi-conducteurs à gap indirect (tel que le silicium) ne sont donc pas adaptés pour former des sources lasers.

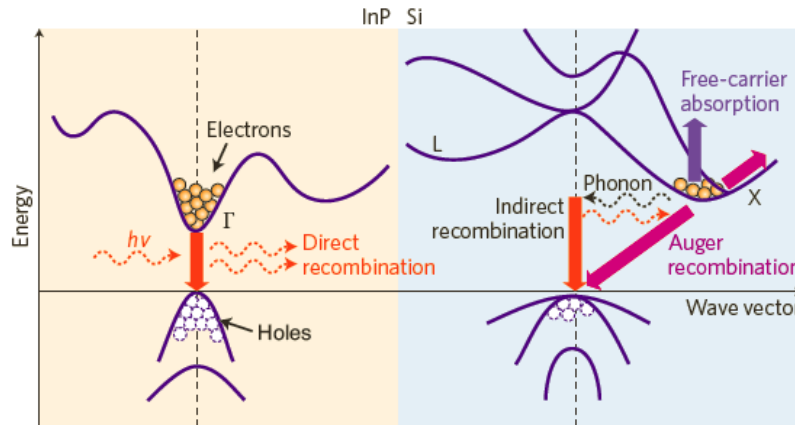


Figure A-2: Diagrammes de bandes d'énergie et principaux phénomènes de transitions interbandes des porteurs de charges dans une structure à bandgap direct tel que l'InP (à gauche), et dans une structure à bandgap indirect tel que le silicium (à droite) [15].

Pour résoudre ce problème, différentes solutions ont été proposées pour obtenir une source laser directement intégrée en photonique sur silicium. Ainsi, malgré ses limitations, un laser silicium basé sur l'effet Raman a été démontré, mais nécessite une pompe optique pour fonctionner [16]. De même, le germanium (déjà intégré sur silicium pour les photo détecteurs) est aussi envisagé comme matériau à gain optique pour former des lasers. En effet, bien que possédant un gap indirect, l'application de contraintes mécaniques peut permettre de modifier suffisamment sa structure de bandes pour le transformer en semi-conducteur à gap direct [22]. Cependant, un laser germanium utilisable pour des applications photoniques (fonctionnement en température ambiante, pompage électrique) reste à démontrer. Pour finir, la croissance sur silicium de matériaux III-V possédant un gap direct (communément utilisés pour former des diodes lasers sur leurs substrats natifs) connaît un intérêt renouvelé de la part de l'industrie microélectronique, qui pourrait les utiliser pour la nouvelle génération de transistors CMOS. Ainsi, une structure laser épitaxiée sur silicium, fonctionnant avec un pompage électrique en régime continu, et émettant à $1.3\mu\text{m}$ a été récemment démontrée [25]. Néanmoins, à cause des différences de paramètres de mailles et de coefficients de dilatations thermiques entre le silicium et les matériaux III-V, de nombreux défauts (tels que des dislocations) peuvent apparaître durant l'épitaxie, et réduire les performances du laser. Pour bloquer la propagation de ces défauts, d'épaisses couches « tampons » placée entre la zone active du laser et le silicium sont nécessaires, ce qui complique le couplage de la lumière dans le circuit photonique silicium (voir **figure A-3(a)**). Une autre structure laser, épitaxiée sans couches tampons, et dont les défauts de croissance sont limités à une épaisseur de 20nm a été récemment démontrée (voir **figure A-3(b)**), mais reste pour le moment limitée à un pompage optique en régime pulsé, à la longueur d'onde de $0.92\mu\text{m}$. Bien que ces résultats soient prometteurs, une grande quantité de travail reste nécessaire pour obtenir des lasers efficaces directement intégrés sur silicium, et dont la lumière peut être couplée dans le circuit photonique.

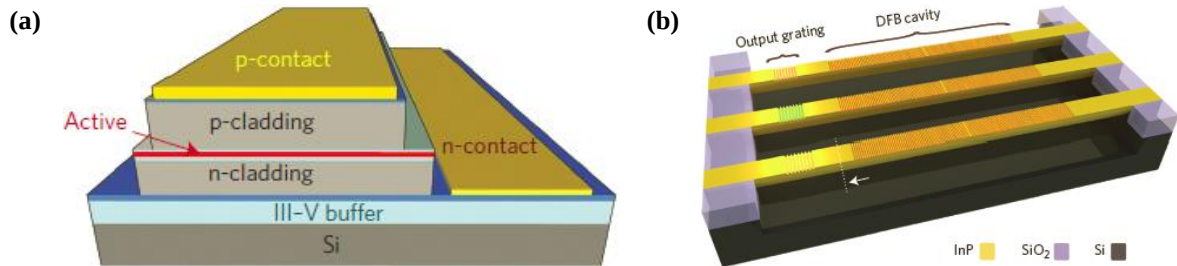


Figure A-3: Vues schématiques de lasers formés par la croissance épitaxiale directe d'InP sur silicium, (a) avec d'épaisses couches tampons III-V [25] et (b) sans couches tampons [26].

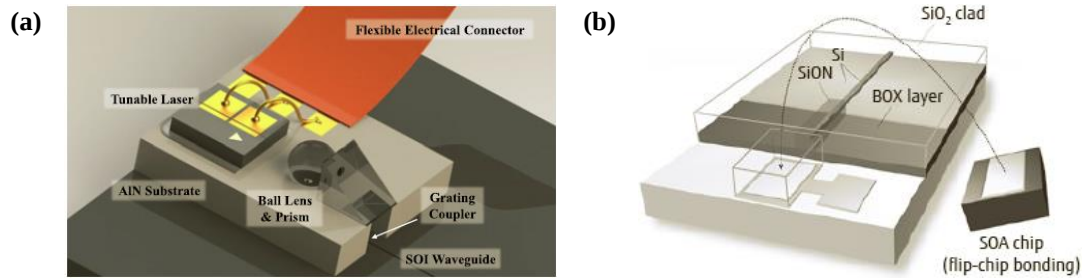


Figure A-4: Schémas de couplages optiques entre une source laser émettant par la tranche et un circuit photonique silicium. (a) Couplage quasi-vertical utilisant des lentilles sphériques et des prismes [29]. (b) Couplage par la tranche utilisant un coupleur épais en SiON et un guide silicium fin [39].

Par conséquent, les systèmes de transmission actuels en photonique sur silicium se basent principalement sur des sources optiques externes au circuit. Celles-ci sont généralement des diodes lasers III-V préfabriquées, et intégrées au-dessus du circuit photonique silicium par collage « flip-chip ». La lumière émise par la tranche du laser peut-être couplée verticalement vers un réseau de couplage dans le circuit photonique à l'aide de lentilles sphériques et de prismes [29] (voir **figure A-4(a)**), ou bien directement dans un guide silicium à l'aide de systèmes de couplage par la tranche, tel que celui illustré sur la **figure A-4(b)**. Bien que cette solution permette de pallier à l'absence d'une source directement intégrée sur silicium, l'intégration externe souffre aussi de plusieurs limitations. Tout d'abord, les pertes optiques causées par le couplage ne sont pas négligeables : entre -1.6dB et -4dB pour l'état de l'art [29], [31]–[36], [38]–[41]. De plus, ces systèmes de couplage sont très sensibles à un désalignement de la source laser. Ainsi, un désalignement dans le plan du circuit photonique supérieur à $\pm 2\mu\text{m}$ résultera en une perte additionnelle supérieure à 1dB [29], [32]–[35], [38]. L'intégration externe requiert donc des techniques de collage par « flip-chip » précises, qui vont réduire le rendement de fabrication, et augmenter le coût et la complexité de la mise en boîtier du circuit photonique complet. Une solution alternative, permettant d'obtenir une source laser intégrée à l'échelle du wafer qui ne souffrirait pas de ces problèmes de couplage optique et d'alignement, reste donc un objectif majeur des recherches en photonique sur silicium.

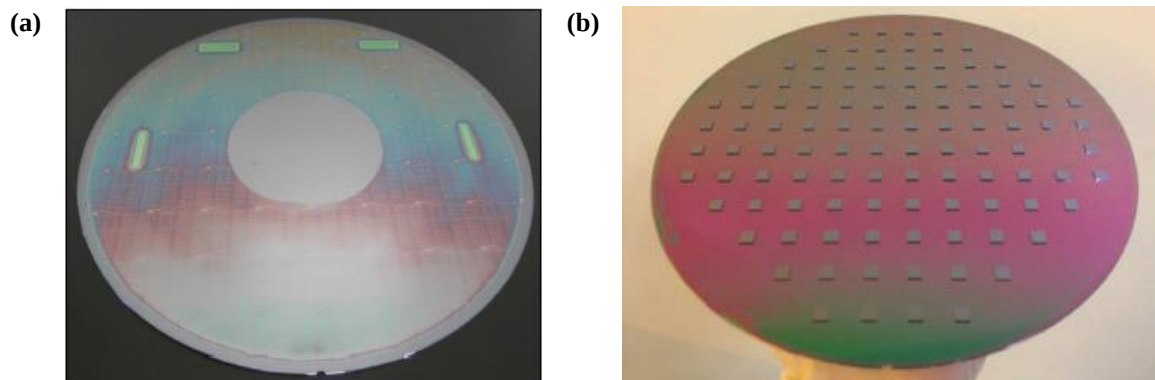


Figure A-5: (a) Plaque de III-V (75mm) collée sur une plaque SOI (200mm). (b) 104 puces de III-V (5mm $\times$ 5mm) collées sur une plaque SOI (200mm) [70].

Cette alternative, située entre l'intégration monolithique et l'intégration externe, est l'intégration hétérogène par collage de matériaux à gap direct (tels que des matériaux III-V) sur des plaques de SOI (pour « silicium on insulator ») pouvant fournir le gain nécessaire pour le laser. Trois types de collages peuvent être trouvés dans la littérature: direct (ou moléculaire) [42]–[58], polymère [59]–[62], et métallique [63]–[68]. Le collage direct reste le plus communément utilisé, et nécessite une couche d'oxyde (dont l'épaisseur peut varier entre quelques nanomètres à une centaine de nanomètres) entre le III-V et le circuit photonique silicium. La surface de l'oxyde doit être aussi propre et plane que possible, afin d'éviter tout défaut de collage. Les matériaux III-V peuvent être collés sur les plaques silicium en tant que plaques complètes (comme sur la **figure A-5(a)**), ou en tant que puces individuelles (comme sur la **figure A-5(b)**). Bien que le collage de plaques soit plus simple à mettre en pratique, celui-ci n'est pas adapté pour de la production. En effet, les plaques de III-V sont généralement limitées à un

diamètre compris entre 75mm et 150mm, alors que celui des plaques de silicium est généralement de 200mm ou 300mm. Les plaques de III-V sont aussi coûteuses, alors que la majorité de la surface collée sera gravée durant la fabrication. Le collage de puces pourrait résoudre ces problèmes, mais leurs rendements de collage, ainsi que la durée totale de l'opération de collage doivent encore être optimisés.

La composition des couches de III-V peut varier, mais comprend généralement une zone de gain, entourée par une épaisse couche ($\approx 1-2\mu m$) d'InP dopée p , et une couche d'InP dopée n utilisées pour former les contacts. Après le collage, différentes opérations de gravure et de métallisations sont effectuées pour obtenir une structure composée d'un guide actif III-V et d'un guide passif silicium, telle que celles visibles sur la **figure A-6**. Dans ces structures, la puissance optique générée par la zone de gain est partagée entre les guides actifs et passifs, selon leurs géométries respectives. Au milieu des guides, la lumière doit être confinée de préférence dans le guide actif pour augmenter le gain modal du laser. A leurs extrémités, celle-ci doit être confinée dans le guide passif pour faciliter le couplage dans le circuit photonique silicium. Les miroirs formant la cavité laser sont généralement présents dans les guides de silicium. Le transfert de puissance entre les deux guides est réalisé grâce à la théorie des modes couplés.

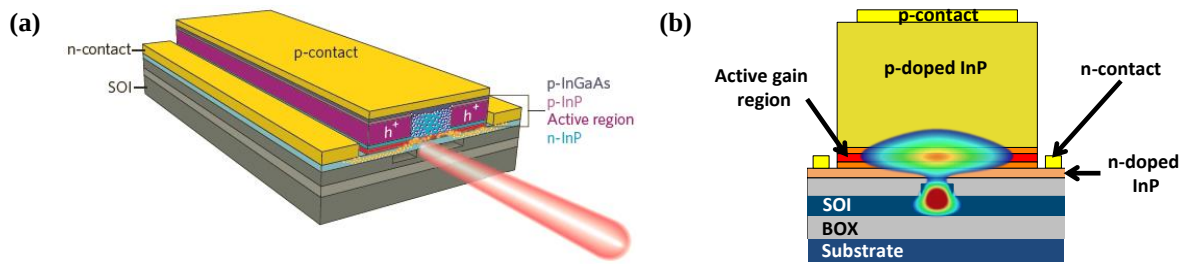


Figure A-6: Vues schématiques d'une structure hybride III-V sur silicium. (a) Vue en trois dimensions [15], et (b) vue en coupe, avec le mode optique distribué entre les guides actifs et passifs.

Pour illustrer cette théorie, prenons le cas de deux guides isolés et sans pertes (1 et 2), chacun parcourus par un mode fondamental avec leurs indices effectifs ($\tilde{n}_1$ et $\tilde{n}_2$) et leurs constantes de propagation respectives ($\tilde{\beta}_{1,2} = 2\pi\tilde{n}_{1,2}/\lambda$, λ étant la longueur d'onde du signal optique). Si les guides sont isolés, l'amplitude du champ électrique de chaque mode (E_1 et E_2) reste alors constante le long des guides. Cependant, s'ils sont suffisamment proches pour que leurs distributions de champs électriques se superposent, E_1 et E_2 deviennent alors dépendant de la distance de propagation (z), et les modes optiques de chaque guide peuvent échanger leur puissance optique. Dans ce cas de figure, il est utile d'utiliser le formalisme des « supermodes » afin d'étudier cette structure à deux guides, et résoudre les équations de couplage qui en découlent [78]–[80]. Ainsi, le champ électrique total est d'abord décrit à l'aide d'un vecteur dans la base des distributions de champ électrique du mode fondamental chaque guide isolé :

$$\hat{E}(x, y, z) = \begin{bmatrix} E_1(z)e^{-j\tilde{\beta}_1 z} \\ E_2(z)e^{-j\tilde{\beta}_2 z} \end{bmatrix} \quad (\text{A.1})$$

Deux solutions aux équations de couplage existent pour le champ électrique total. Ces solutions correspondent à des supermodes pairs (E_e) et impairs (E_o), caractérisés chacun par une constante de propagation différente (respectivement $\tilde{\beta}_e$ et $\tilde{\beta}_o$). L'amplitude du champ électrique de ces supermodes dans chaque guide dépend de la différence entre les constantes de propagation des modes fondamentaux chaque guide isolé ($\delta = (\tilde{\beta}_1 - \tilde{\beta}_2)/2$), et d'une constante de couplage (κ), représentant la force de couplage entre les deux guides. Cette force de couplage augmente avec la superposition des modes fondamentaux de chaque guide isolé. Dans certains cas particuliers, l'amplitude du champ électrique des supermodes pairs et impairs peut être exprimée de façon simplifiée. Ces différents cas particuliers sont illustrés sur la **figure A-7**, et permettent d'utiliser deux types de couplages : directionnel et adiabatique.

Le cas du couplage directionnel correspond à la **figure A-7(b)**, qui représente la condition d'accord de phase entre les deux guides ($\delta = 0$). Dans ce cas, chaque supermode est distribué de manière égale dans chaque guide, et bien que leurs amplitudes de champs respectives restent constantes, leurs constantes de propagation sont différentes. Ainsi, au fur et à mesure qu'ils se propageront dans les guides, ils interféreront alternativement de

manière destructive ou constructive dans chaque guide, permettant ainsi un échange périodique de puissance entre les deux guides. Cependant, ce type de couplage nécessite un accord de phase presque parfait pour fonctionner, et n'est donc pas adapté pour un couplage entre deux guides basés sur des matériaux différents, comme dans le cas d'une structure hybride.

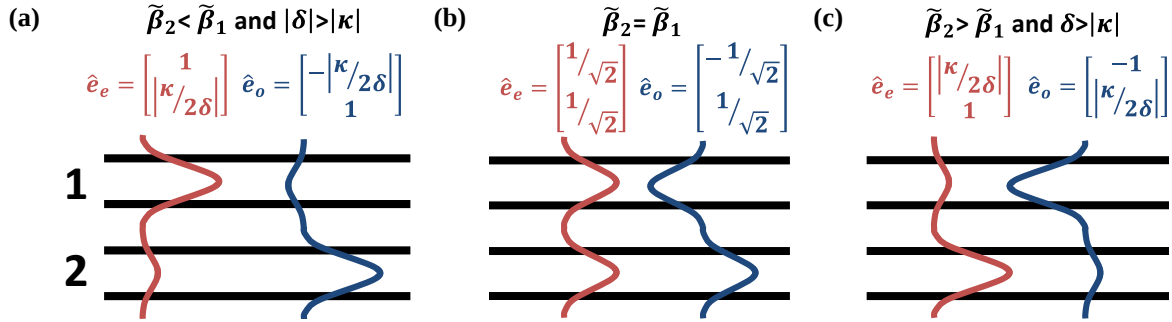


Figure A-7: Cas particuliers du couplage entre deux guides: (a) le supermode pair/impair est principalement confiné dans le guide 1/2, (b) la puissance optique de chaque supermode est équitablement répartie dans chaque guide (adaptation de phase), et (c) le supermode impair/pair est principalement confiné dans le guide 2/1.

Le couplage adiabatique se base sur la transformation des supermodes, illustrée sur les **figures A-7(a) et (c)**. En contrôlant la différence entre les constantes de propagation (ou indices effectifs) des modes fondamentaux des guides isolés, la puissance optique contenue dans les supermodes peut être transférée d'un guide à un autre. Ce contrôle peut être obtenu en modifiant la largeur des guides, c'est-à-dire en formant des « tapers » dans les guides. Bien qu'un transfert total (100%) de la puissance optique soit impossible et qu'ils soient plus longs que les coupleurs directionnels, les coupleurs adiabatiques sont plus robustes et adaptés pour des structures hybrides.

Dans le cas d'un laser hybride, seul le supermode avec le meilleur compromis entre confinement dans la zone active et contre-réaction des miroirs dans la cavité. En général, le supermode pair est celui choisi lors de la conception de la structure hybride, comme illustré sur la **figure A-8**.

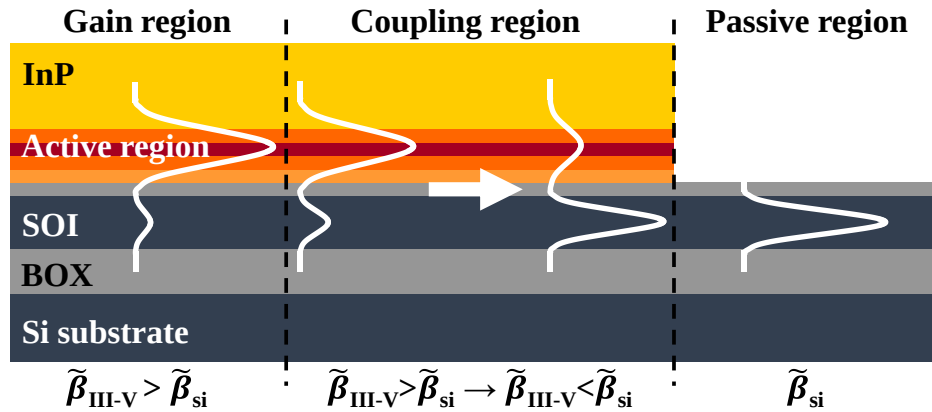


Figure A-8: Vue schématique du transfert de puissance entre les guides actifs et passifs, réalisé par couplage adiabatique.

Depuis la première démonstration du premier laser hybride III-V sur silicium à pompage électrique en régime continu [42], de nombreuses démonstrations de lasers hybrides ont été réalisées avec différents types de cavités : Fabry-Pérot[44]–[46], [59], circuits [46], [47], miroirs de Bragg (ou DBR, équivalent de « Distributed Bragg Reflectors ») [48], [51]–[53], [67], contre-réaction distribuée (ou DFB, équivalent de « Distributed Feed-Back »)[49]–[51], [60], micro-disques [54], et micro-anneaux [55]. Bien que n'étant pas directement intégrés sur silicium, ces sources lasers présentent plusieurs avantages si on les compare à des sources laser externes. Ainsi, les seules pertes de couplages entre le laser et le circuit silicium proviennent du taper adiabatique, et peuvent être largement réduites [83]. De plus, le collage de matériaux III-V ne nécessite pas d'alignement précis, puisque celui-ci est réalisé durant les étapes de lithographie des guides III-V. Par ailleurs, les sources étant intégrées à

l'échelle des plaques de silicium, la mise en boîtier du circuit photonique s'en retrouve simplifiée, et son coût devrait ainsi diminuer. Finalement, l'intégration hétérogène de matériaux III-V offre de nouvelles possibilités dans la conception des lasers.

Malgré tous ces avantages, les démonstrations de sources lasers III-V sur silicium sont généralement celles de composants isolés, ou au mieux intégrés avec des composants siliciums passifs. L'absence de démonstration de systèmes plus complexes peut s'expliquer par les problématiques d'intégration de matériaux III-V sur les plateformes de fabrication de photonique sur silicium. En effet, bien que les matériaux III-V présentent un intérêt pour la microélectronique, ceux-ci restent des matériaux « exotiques ». De même, les couches de contacts électriques sur III-V comprennent généralement de l'or, qui est interdit dans les plateformes de fabrication CMOS. Un autre problème concerne le guide de silicium situé sous le III-V, qui doit être suffisamment épaisse (généralement $> 400nm$) pour permettre le couplage optique entre le III-V et le silicium, alors que les plateformes standards de fabrication en photonique sur silicium utilisent des couches plus fines ($< 300nm$). Pour finir, il existe une incompatibilité entre l'empilement de couches III-V, qui mesure entre 2 et $3\mu m$, et les différents niveaux de métallisation d'un circuit photonique standard, dont la hauteur est du même ordre de grandeur. Le laser, ainsi que chaque niveau de métal nécessitant chacun une étape de planarisation, la configuration de la **figure A-9** est donc impossible à réaliser en l'état.

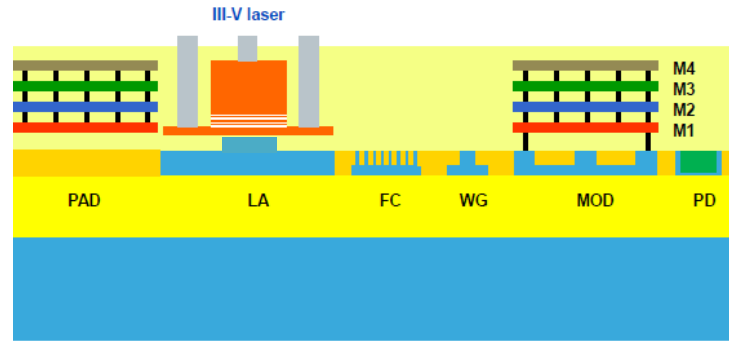


Figure A-9: Vue schématique d'un circuit photonique sur silicium, comprenant différents types de composants actifs et passifs, illustrant les problèmes d'intégration d'une source laser en termes d'épaisseur de couche SOI, et de niveaux de métallisation standards. Cette co-intégration est impossible dans ces conditions.

A cause des difficultés liées à l'intégration d'une source laser au sein d'un circuit photonique sur silicium, il existe seulement deux démonstrations de tels transmetteurs intégrés à haut débit. L'une de ces démonstrations a été réalisée par Alduino et al. (voir **figure A-10(a)**), avec un transmetteur basé sur quatre lasers DBR hybrides III-V sur silicium émettant chacun à une longueur d'onde fixée autour de $1.3\mu m$, et co-intégrés avec quatre modulateurs silicium fonctionnant chacun à un débit de $12.5Gb/s$ [84], [85]. La seconde a été proposée par Duan et al. (voir **figure A-10(b)**), avec un transmetteur intégrant un unique laser accordable en longueur d'onde autour de $1.55\mu m$ (jusqu'à $9nm$), et un modulateur silicium limité à un débit de $10Gb/s$ [86]. Malgré cette limitation en débit de modulation, l'accordabilité en longueur d'onde demeure essentielle pour réduire l'espacement entre les différentes longueurs d'ondes (ou canaux) du transmetteur.

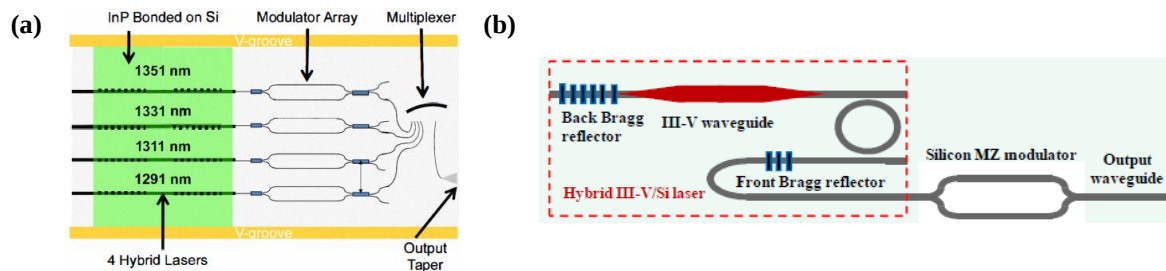


Figure A-10: Vues schématiques de transmetteurs à haut débit, co-intégrant des lasers hybrides III-V sur silicium et des modulateurs silicium, avec : (a) quatre lasers à longueur d'onde fixe [85], et (b) un unique laser accordable en longueur d'onde [86].

Une autre approche pour un transmetteur basé sur des sources lasers hybrides III-V sur silicium est d'utiliser les matériaux III-V apportés par hétéro-intégration pour former les différents composants actifs du circuit photonique (modulateurs et photo-détecteurs), et de n'utiliser le SOI que pour former les composants passifs. Cette approche a été utilisée par Ramaswamy et al. (voir **figure A-11**), afin de réaliser un transmetteur basé sur quatre lasers hybrides III-V sur silicium, émettant à quatre longueurs d'ondes fixées autour de $1.3\mu\text{m}$, et co-intégrés avec quatre modulateurs III-V sur silicium fonctionnant chacun à un débit de 28Gb/s [87]. Bien que cette solution permette de résoudre plusieurs problèmes d'intégration liés au laser, elle n'est pas triviale pour autant, comme détaillé dans la section suivante.

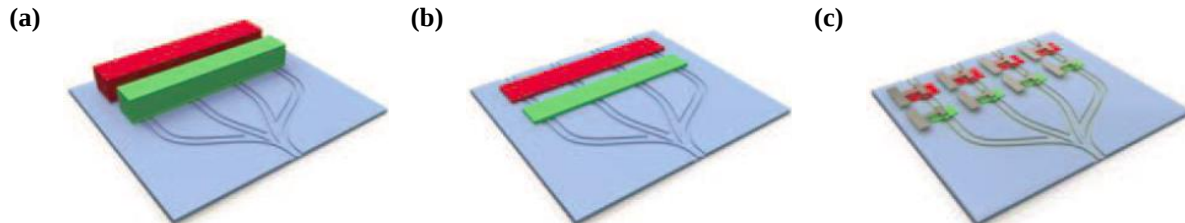


Figure A-11: Transmetteur hybride III-V sur silicium, où le III-V est utilisé pour former les différents composants actifs du circuit photonique. Vues schématisques du transmetteur (a) après collage, (b) après retrait du substrat, et (c) après les procédés de gravure et de métallisation. Images provenant de [88].

Puisque les transmetteurs intégrés ne sont pas encore répandus en photonique sur silicium, le but de ce travail de thèse est de proposer notre propre vision pour un transmetteur hybride III-V sur silicium à haut débit, aussi compatible que possible avec les plateformes de fabrication silicium actuelles. L'application choisie pour ce transmetteur étant la transmission d'information dans les centres de données, ses spécifications ont été basées sur celles de la norme « 100GBASE-LR4 » (provenant du standard IEEE-802.3ba). Ainsi, les transmetteurs conçus devront transmettre des informations jusqu'à 10km, en utilisant quatre longueurs d'onde autour de $1.3\mu\text{m}$, avec un espacement de 4.5nm , chacune modulée à un débit de 25Gb/s . Les lasers intégrés dans les transmetteurs doivent émettre à une seule longueur d'onde (monomodes), c'est-à-dire avec un SMSR (« Side-Mode Suppression Ratio ») supérieur à 30dB . Par ailleurs, l'OMA (« optical modulation amplitude », défini par la différence en mW entre les niveaux haut et bas du signal optique modulé) à chaque longueur d'onde devrait être supérieure à -1.3dBm , afin de limiter les erreurs de lecture pendant la détection du signal. Ces différents aspects sont traités dans les sections A-2 (pour le modulateur) et A-3 (pour le transmetteur complet) de cette annexe. Un autre objectif de ce travail de thèse est d'améliorer l'intégration des lasers hybrides avec le reste du circuit photonique sur silicium, et en particulier la problématique de l'épaisseur de SOI nécessaire pour coupler le signal provenant du laser dans le circuit. Cet aspect est abordé dans la section A-4, suivi de la section A-5 qui contient les différentes conclusions pouvant être tirées de ce travail, ainsi que ses conclusions.

A-2. Modulateur Silicium

La première étape de ce travail de thèse fut de concevoir le modulateur à intégrer dans le transmetteur. Afin de moduler un signal optique, deux options sont possibles : 1) moduler le coefficient d'absorption du matériau constituant le guide d'onde pour contrôler directement l'amplitude du signal optique, ou 2) moduler son indice de réfraction pour contrôler la phase du signal optique, et utiliser une structure interférométrique pour modifier son amplitude. Les modulateurs basés sur le premier effet sont appelés des EAMs (pour « electroabsorption modulators »), et ceux basés sur le second des EOMs (pour « electrooptic modulators »).

Les EAMs sont généralement basés sur l'effet Franz-Keldysh, qui consiste à modifier l'absorption d'un matériau à l'aide d'un champ électrique. Cependant, cet effet est bien plus efficace sur les matériaux à gap direct, et n'est donc pas adapté pour le silicium. De plus, l'effet de modulation fonctionne essentiellement au bord de la bande d'absorption [92]. Une solution consiste donc à utiliser une structure hybride III-V sur silicium afin de former des EAMs. Ceux-ci sont généralement basés sur le « quantum confined Stark effect » (ou QCSE), qui est l'équivalent de l'effet Franz-Keldysh pour des puits quantiques. Bien que plus efficaces, le fonctionnement des modulateurs basés sur le QCSE est encore plus restreint autour du bord de la bande d'absorption, ce qui limite leur bande passante optique, les rendant aussi plus sensibles aux variations de température [3]. Des modulateurs

hybrides III-V sur silicium très efficaces (tels que ceux illustrés sur la **figure A-12**) ont ainsi été démontrés, avec des taux d'extinction (ratio entre les puissances optiques correspondant aux bits "0" et "1") proche de $10dB$ et des pertes d'insertion de $5dB$, fonctionnant à des débits de $50Gb/s$, avec des tensions de $2V$ [96], [97].

Bien qu'une approche basée sur l'utilisation de matériaux III-V semble appropriée pour former à la fois les sources lasers et les modulateurs d'un transmetteur hybride III-V sur silicium, elle est aussi relativement compliquée à mettre en œuvre. Pour commencer, le nombre de puits quantiques utilisé pour le modulateur est généralement plus élevé que pour le laser, afin de réduire sa capacitance [96]. De plus, la composition de ces puits quantiques doit aussi être différente, afin de décaler le spectre d'absorption du modulateur vers des longueurs d'ondes plus faible. Ainsi, la longueur d'onde émise par le laser se retrouvera à la bordure du spectre d'absorption du modulateur, où l'effet de modulation est le plus élevé. Utiliser le même empilement de matériaux III-V pour former à la fois les lasers et les modulateurs résulte donc en un compromis dans les performances de chaque composant. Il est donc nécessaire d'employer des techniques telles que le « quantum-well intermixing » (ou QWI), afin de modifier localement la structure de bande de chaque composant, mais cette méthode reste limitée pour former des transmetteurs complets, à cause de la difficulté d'ajuster les structures de bande [100]–[102]. Une autre solution est d'utiliser plusieurs puces de III-V, chacune avec un empilement III-V différent pour chaque composant. Cependant, cette approche requiert un collage proche de chaque puce pour former des circuits compacts, ce qui accroît la complexité de fabrication. Cette solution ne fut donc pas retenue pour former les modulateurs.

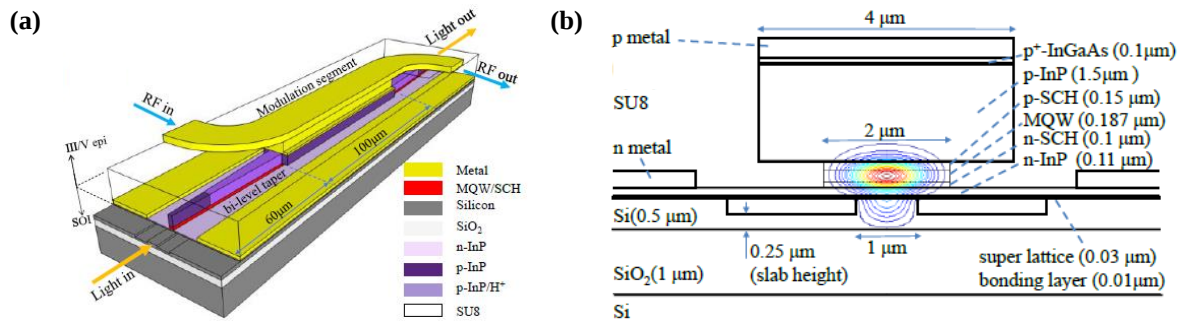


Figure A-12: Vues schématiques (a) de dessus et (b) transverse d'un EAM III-V sur silicium. Images tirées de [97].

Plusieurs effets existent pour modifier l'indice de réfraction d'un matériau et former des EOMs. Cependant, les effets électro-optiques tels que l'effet Pockels ou Kerr, sont inexistantes ou trop faibles dans le silicium. L'effet thermo-optique est très efficace dans le silicium, mais reste trop lent (de l'ordre de la microseconde) pour former des modulateurs aux débits ciblés. Cet effet reste néanmoins utile pour induire une différence de phase statique dans les modulateurs, indispensable à leur bon fonctionnement.

La solution la plus utilisée pour former des modulateurs en silicium est donc d'introduire des porteurs libres dans le guide silicium par implantation ionique, afin de modifier son indice de réfraction. Ainsi, en contrôlant la distribution spatiale des porteurs dans le guide d'onde, il devient possible de moduler le signal optique. Néanmoins, l'ajout de ces porteurs va aussi augmenter l'absorption dans les zones implantées, et donc augmenter les pertes optiques du modulateur. Trois mécanismes sont généralement utilisés pour contrôler cette distribution de porteurs :

1. **L'injection de porteurs :** Les structures basées sur l'injection de porteurs utilisent une jonction *p-i-n* implantée dans un guide d'onde. Quand celle-ci est polarisée en direct, les porteurs sont injectés dans le guide, modifiant ainsi son indice effectif (voir **figure A-13(a)**). Bien qu'étant le plus efficace des trois mécanismes [107], il est aussi limité en terme de bande passante électro-optique par le temps de vie des porteurs minoritaires (à quelques centaines de mégahertz [105]). Ainsi, les modulateurs basés sur l'injection de porteurs nécessitent l'utilisation de signaux d'entrée distordus pour fonctionner à des débits de l'ordre de 25 ou $50Gb/s$ (on parle alors de « pre-emphasis ») [108]–[110], qui augmentent alors la complexité du circuit de commande.

2. **L'accumulation de porteurs :** Les structures basées sur l'accumulation de porteurs requièrent une fine couche isolante située entre deux couches semi-conductrices dopées *p* et *n*, à l'intérieur d'un guide d'onde. Lorsque la structure est polarisée en directe, les porteurs s'accumulent de chaque côté de la barrière isolante et

modifient l'indice effectif du guide (voir **figure A-13(b)**). Bien que moins efficace que l'injection, l'accumulation n'est limitée en terme de rapidité que par la capacitance de la couche isolante, et possède une bande passante électro-optique plus élevée [105]. Cependant, ces composants sont aussi plus complexes à fabriquer, et ceux démontrés jusqu'à présent avec un débit supérieur à 10Gb/s sont rares [112], [114], [115].

3. La déplétion de porteurs : Les structures basées sur la déplétion de porteurs utilisent une jonction $p-n$ située dans un guide d'onde. Quand celle-ci est polarisée en inverse, les porteurs présents sont retirés du guide, modifiant ainsi son indice effectif (voir **figure A-13(c)**). Bien qu'étant le moins efficace des trois mécanismes [107], il est aussi beaucoup moins limité en terme de bande passante électro-optique, grâce à sa faible capacitance de jonction. Sa simplicité de fabrication est aussi l'un des facteurs qui nous amené à choisir ce mécanisme pour notre transmetteur intégré.

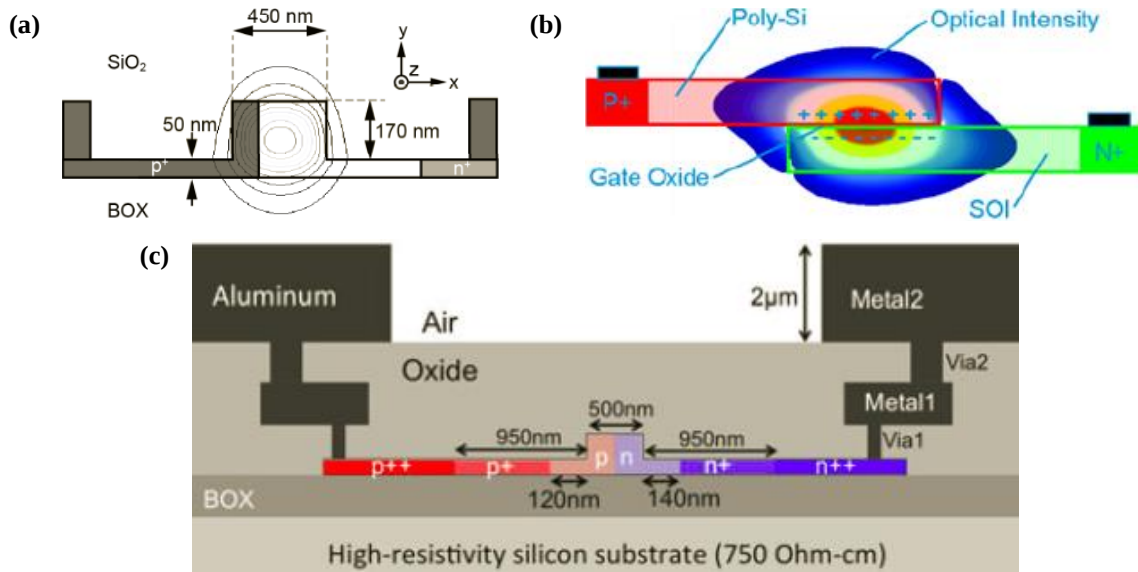


Figure A-13: Vues transversales schématiques de modulateurs basés sur (a) l'injection [110], (b) l'accumulation [112], et (c) la déplétion [127] de porteurs libres.

L'étape suivante est de choisir la structure à utiliser pour former le modulateur. Deux types de structures sont généralement dans la littérature : les anneaux (voir **figure A-14(a)**) ou les interféromètres de Mach-Zehnder (voir **figure A-14(b)**). Bien qu'étant compacts et pouvant utiliser une faible tension de fonctionnement, les modulateurs en anneaux sont des structures résonnantes, avec des bandes passantes optiques réduites et sont donc extrêmement sensibles aux variations de températures et de fabrication. D'un autre côté, les interféromètres de Mach-Zehnder possèdent une large bande passante optique, et sont donc intrinsèquement robustes face aux variations de fabrication et de températures. Ils furent donc choisis pour être intégrés dans le transmetteur. Cependant, à cause de l'efficacité réduite de la déplétion de porteur sur la modulation de l'indice effectif, les modulateurs Mach-Zehnder (ou MMZ), utilisent des longueurs de l'ordre du millimètre. Ainsi, cela risque d'augmenter leurs pertes optiques et de limiter leur rapidité de fonctionnement à cause d'une capacitance trop élevée, s'ils ne sont pas conçus avec soin.

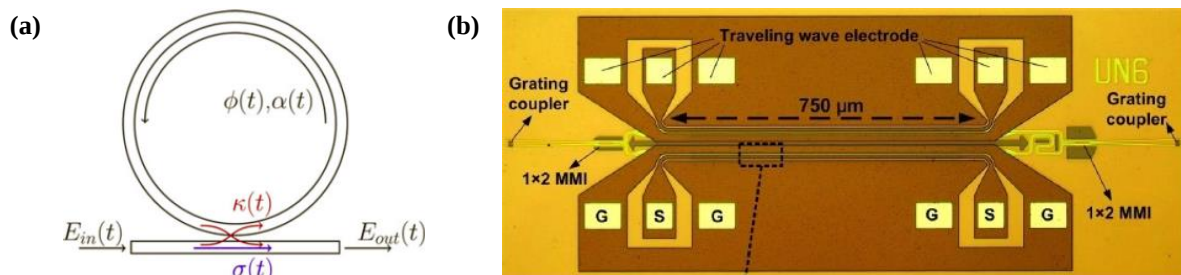


Figure A-14: (a) Vue schématique de dessus d'un modulateur en anneau [132]. (b) Vue de dessus d'un MMZ silicium [123].

Les modulateurs présentés dans ce document sont basés sur une couche de SOI de 300nm , avec un « slab » de 150nm , afin d'être compatible avec les plateformes de fabrication de STMicroelectronics et du CEA-LETI. Comme ils seront intégrés à terme avec un laser hybride III-V sur silicium, dont les électrodes sont réalisés par la méthode du « lift-off », un unique niveau de métal sera accessible pour former les électrodes du modulateur. Pour finir, afin de pouvoir éventuellement utiliser un circuit électronique pour contrôler le composant, ses tensions de fonctionnement doivent être limitées à 2.5V .

Une vue transverse de l'un des bras du MMZ est illustrée sur la **figure A-15**. Celui-ci comporte une jonction p - n latérale (choisie pour sa simplicité de fabrication), ainsi que deux zones de contact avec un dopage élevé pour réduire leurs résistances. Bien que les zones de contact doivent être suffisamment proches pour réduire la résistance d'accès du modulateur, elles risquent d'augmenter les pertes optiques si trop proches du mode optique. Des simulations par éléments finis ont permis de déterminer que la distance séparant les bords du guide optique aux zones de contact doit être supérieure à 500nm , et pour réduire les risques dus à la fabrication, cette distance a finalement été fixée à 800nm . Par ailleurs, d'après les relations semi-empiriques de Soref [92], les trous sont plus efficaces que les électrons pour modifier l'indice de réfraction du silicium, et apportent moins d'absorption. La jonction a donc été décalée de 100nm par rapport au centre du guide, afin de favoriser l'effet des trous.

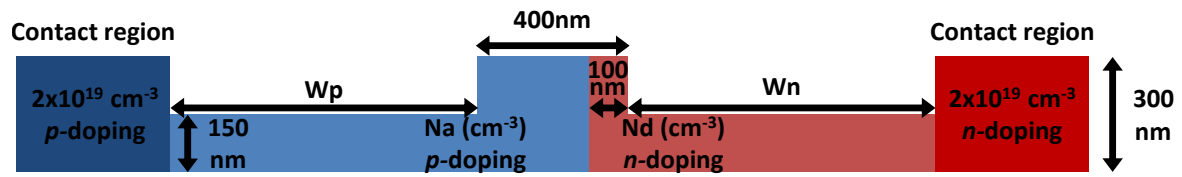


Figure A-15: Vue transversale schématique de l'un des bras du MMZ silicium.

L'étape suivante est de trouver la bonne combinaison de concentrations de dopages pour former la jonction p - n du modulateur, afin d'obtenir le meilleur compromis entre l'efficacité de modulation de l'indice effectif du guide optique et ses pertes de propagation. Un modèle « réaliste » de la jonction (illustré sur la **figure A-16(a)**), prenant en compte les différentes étapes de fabrication de la jonction (recuits, implantations, oxydations), a ainsi été simulé pour différentes concentrations de dopages. Trois implantations successives à différentes énergies sont réalisées de chaque côté de la jonction pour obtenir un profil vertical de dopage uniforme. Le résultat de ces simulations est synthétisé sur la **figure A-16(b)**, avec l'efficacité de modulation de la jonction polarisée en inverse à une tension de 2.5V (représenté par le décalage de phase du signal optique), et les pertes de propagation liées à la jonction polarisée à 0V pour différentes combinaisons de dopages.

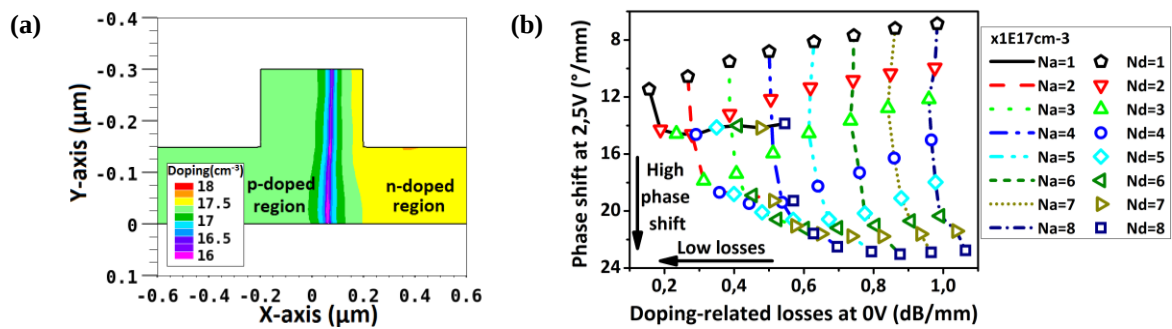


Figure A-16: (a) Différence des concentrations des dopages n et p en valeur absolue pour un profil de dopage « réaliste » prenant en compte les étapes de fabrication de la jonction. (b) Décalage de phase du signal optique à une tension de -2.5V , par rapport aux pertes de propagation à 0V pour différentes concentrations de jonctions. Chaque point représente une combinaison différente de concentrations d'atomes accepteurs (Na) et donneurs (Nd).

Ainsi, bien qu'augmenter les dopages des régions p et n de la jonction permette d'augmenter l'efficacité de la jonction, cela augmente aussi très largement les pertes optiques. Cependant, il semble préférable d'augmenter le dopage de la zone n (atomes donneurs) plutôt que celui de la zone p (atomes accepteurs). Finalement, les concentrations des zones p et n sont respectivement fixées à $4 \times 10^{17}\text{cm}^{-3}$ et $6 \times 10^{17}\text{cm}^{-3}$. Cette jonction

devrait fournir un décalage de phase de $21.5^\circ/mm$ à $-2.5V$ (équivalent à un $V_\pi L_\pi$ de $2.1V.cm$) avec des pertes liées aux dopages de la jonction de $0.63dB/mm$ à $0V$. Les performances attendues de cette jonction, ainsi que sa capacitance, sont résumées sur la **figure A-17**. Trois longueurs de zone active basée sur cette jonction ont été choisies pour les modulateurs afin de trouver le meilleur compromis entre efficacité et pertes: $2mm$, $4mm$ et $6mm$.

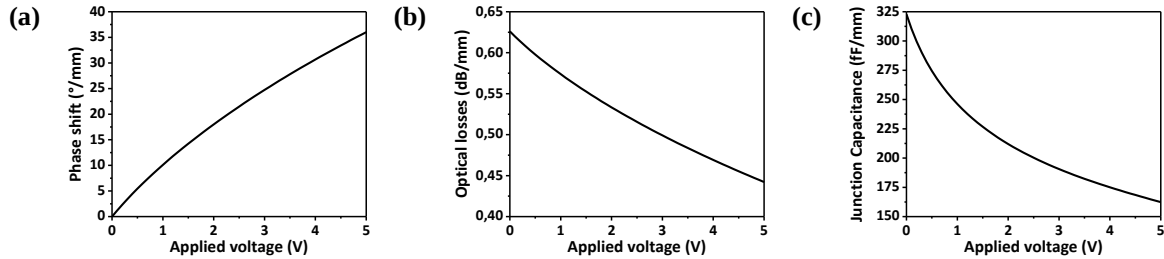


Figure A-17: (a) Décalage de phase, (b) pertes liées aux dopages de la jonction, et (c) capacité de jonction en fonction de la tension inverse appliquée sur la jonction choisie pour le modulateur.

La dernière étape de la conception du modulateur consiste à étudier son comportement à hautes fréquences, pour moduler des signaux à haut débit. Celui-ci étant basé sur la déplétion de porteurs, il n'est pas limité par le temps de transit des porteurs de charges (évalué à quelques picosecondes), mais par sa capacitance, qui est relativement élevée pour les longueurs de zones actives choisies. Puisque l'amplitude du signal de commande du modulateur a été fixée à $2.5V$, son point de fonctionnement a été fixé à $-1.25V$. A cette tension, la capacitance de la jonction est estimée à $C_j = 235fF/mm$. Ainsi, si le modulateur est considéré comme une simple capacité, et contrôlée par un générateur avec une impédance $Z_G = 50\Omega$, sa fréquence de coupure RC (calculée par $1/2\pi Z_G C_j$) sera inférieure à $7GHz$, et ce même pour la plus petite longueur de zone active ($2mm$). Cette valeur de bande passante ne tenant pas compte des éventuels composants parasites est donc trop faible pour atteindre le débit ciblé de $25Gb/s$.

Afin d'augmenter la bande passante du modulateur, il est nécessaire de le considérer comme une ligne de transmission, dont la capacitance est distribuée le long de la zone active. Dans ce cas, le modulateur est dit à ondes progressives (ou « travelling-wave »), et nécessite une étude spécifique de ses électrodes afin d'améliorer la propagation du signal électrique. La réponse fréquentielle du modulateur est déterminée par trois paramètres, qui dépendent de sa géométrie et de la fréquence du signal RF : son impédance caractéristique, l'indice effectif de l'onde RF, et les pertes de propagation de l'onde RF. Afin de maximiser la bande passante du modulateur il est nécessaire de satisfaire au mieux les conditions suivantes :

- 1) L'indice effectif de l'onde RF doit être aussi proche que possible de l'indice de groupe de l'onde optique (évalué à ≈ 4 avec notre géométrie de guide d'onde), afin que les deux ondes se propagent approximativement à la même vitesse.
- 2) L'impédance caractéristique du modulateur doit être aussi proche que possible de celles du générateur et de la charge (généralement fixées à 50Ω), afin de réduire les réflexions sur la ligne de transmission.
- 3) Les pertes de propagation sur la ligne doivent être réduites au maximum.

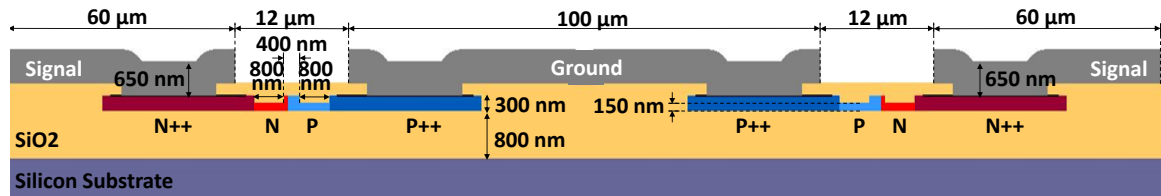


Figure A-18: Vue schématique transverse du modulateur Mach-Zehnder avec ses dimensions définitives.

Ces conditions ne pouvant pas être toutes satisfaites en même temps, un compromis fut trouvé après plusieurs simulations de différentes géométries d'électrodes, à différentes fréquences. Le résultat est illustré sur la **figure A-18**. La structure définitive choisie est à base d'électrodes d'AlCu coplanaires, avec une configuration

signal-masse-signal, où les jonctions $p-n$ de chaque bras du modulateur sont symétriques. Ainsi, il n'y a pas d'asymétrie entre les électrodes, qui pourrait exciter des modes RF non voulus et créer des bosses dans la réponse électro-optique du modulateur. Cette configuration permet aussi de contrôler séparément chaque bras du modulateur, afin d'améliorer ses performances.

Les paramètres caractéristiques de cette géométrie de modulateur obtenus par simulation sont visibles sur la **figure A-19**. Les simulations ont été effectuées pour des tensions inverses de 0V et 1.25V appliquées sur l'un des bras, et prennent aussi en compte la résistivité du substrat de silicium : celui-ci peut être de résistivité standard ($10\Omega.cm$), ou à haute résistivité ($750\Omega.cm$). Au point de fonctionnement du modulateur (1.25V), l'indice effectif de l'onde RF est proche de 4 (les ondes optiques et RF devraient donc se propager à la même vitesse) alors que l'impédance caractéristique est proche de 35Ω (au lieu des 50Ω visés). En effet, il est impossible dans cette configuration de satisfaire les deux conditions : l'indice effectif de l'onde RF est proportionnel à la racine carrée de la capacitance totale entre les électrodes (qui est principalement dépendante de la capacitance de la jonction $p-n$), alors que l'impédance caractéristique est proportionnelle à l'inverse de la racine carrée de cette capacitance totale [90]. Ainsi, lorsqu'une tension inverse de 1.25V est appliquée sur la jonction, la capacitance totale entre les électrodes diminue, et l'impédance caractéristique augmente, alors que l'indice effectif RF décroît, de même, que les pertes RF. L'utilisation d'un substrat à haute résistivité permet aussi de réduire drastiquement les pertes RF de la ligne, et est donc essentielle.

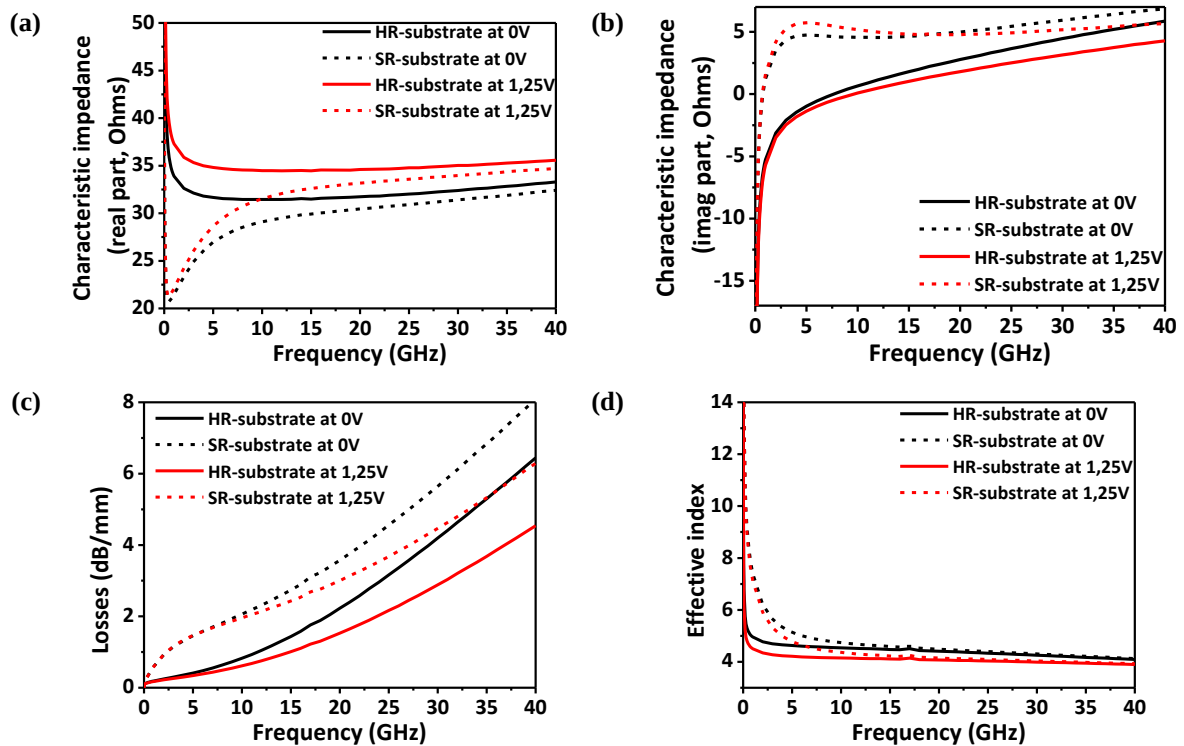


Figure A-19: Simulations des parties (a) réelle et (b) imaginaire de l'impédance caractéristique, (c) des pertes RF et (d) de l'indice effectif de la ligne. Traits continus: simulations avec un substrat à haute résistivité ($750\Omega.cm$). Traits pointillés: simulations avec un substrat standard ($10\Omega.cm$).

En utilisant ces paramètres caractéristiques, il est possible d'évaluer les réponses électro-optiques du modulateur pour différentes longueurs de zones actives. Ces réponses électro-optiques sont visibles sur la **figure A-20**. Ainsi, les fréquences de coupures à $-3dB$ (qui correspondent à une atténuation de 50% de l'amplitude du signal en fin de ligne) sont respectivement estimées à $35.7GHz$, $24.3GHz$ et $19.3GHz$ pour des zones actives de $2mm$, $4mm$ et $6mm$, avec une tension indirecte de 1.25V. Dans chaque cas, l'utilisation d'un substrat à haute résistivité ou d'une tension de polarisation inverse plus élevée permet d'augmenter la bande passante électro-optique du composant, grâce à la réduction des pertes RF de la ligne. Bien que ces valeurs soient suffisantes pour le débit visé de $25Gb/s$, il est encore possible de les améliorer. Par exemple, une valeur plus faible de

capacitance de jonction (obtenue en diminuant les dopages de la jonction) permettrait de réduire la capacitance totale, et donc de trouver un meilleur compromis entre l'impédance caractéristique et indice effectif RF.

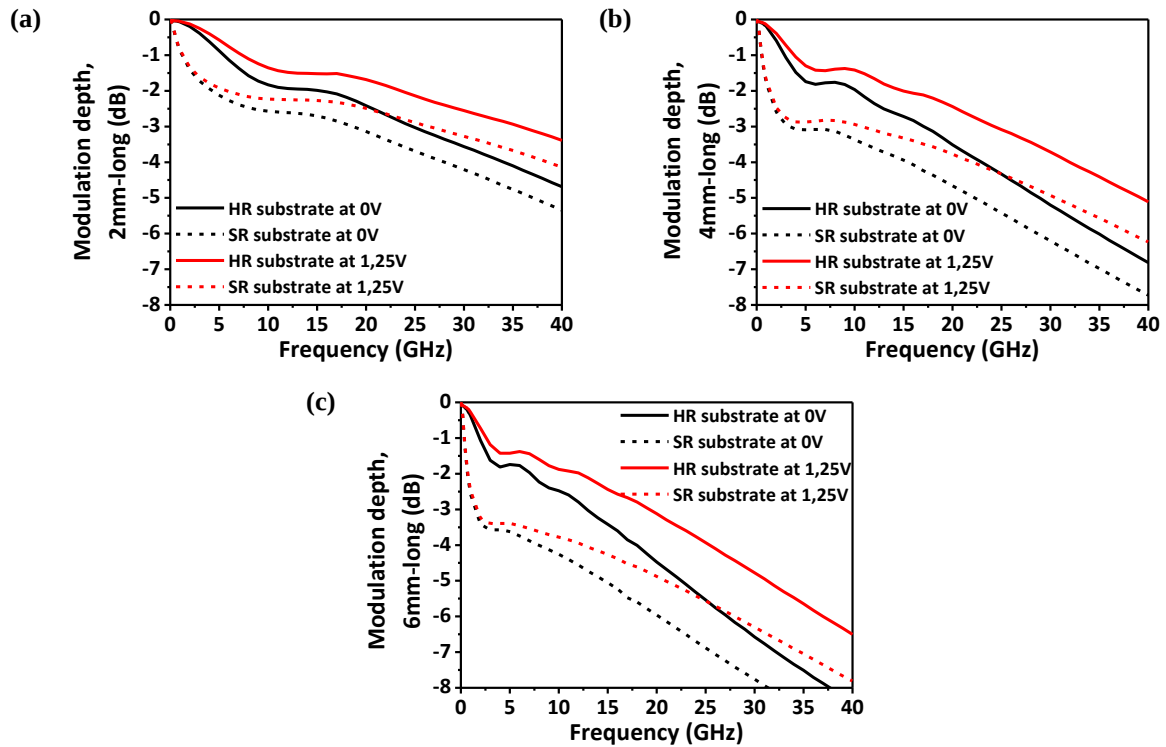


Figure A-20: Simulations de la réponse électro-optique du modulateur pour différentes longueurs de zones actives : (a) 2mm, (b) 4mm, and (c) 6mm. Traits continus: simulations avec un substrat à haute résistivité (750 $\Omega \cdot \text{cm}$). Traits pointillés: simulations avec un substrat standard (10 $\Omega \cdot \text{cm}$).

Les masques de la structure complète du modulateur silicium sont présentés sur la **figure A-21**. La lumière est couplée dans et en dehors de la structure à l'aide de réseaux de couplages vers des fibres optiques monomodes. La lumière est séparée et regroupée à l'aide d'interféromètres multimodes. La région active comprend deux bras implantés avec la jonction p - n , et contrôlés par les électrodes à ondes progressives. Une longueur de 100 μm a été ajoutée sur l'un des deux bras, afin de rendre la structure asymétrique, et simplifier ainsi sa caractérisation électro-optique. Une section passive de 1mm est aussi présente avant la zone active. Si une chauffette est présente au-dessus de cette section, il est alors possible de contrôler la différence de phase statique entre les deux bras, pour adapter le point de fonctionnement du modulateur à la longueur d'onde de la source laser.

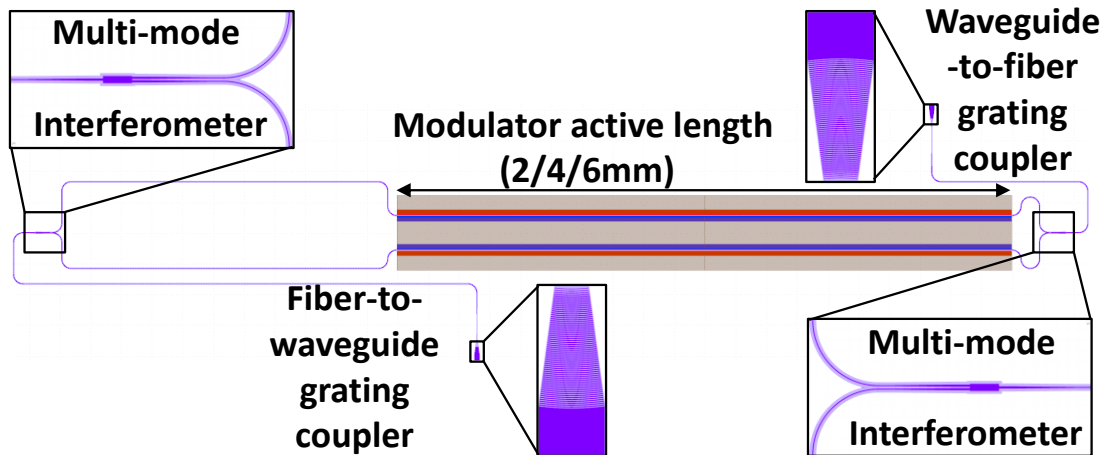


Figure A-21: Masques de la structure complète du modulateur.

Afin de valider la conception de ces modulateurs, des échantillons ont été fabriqués en utilisant les étapes décrites sur la **figure A-22**. La fabrication débute sur des substrats SOI de 200mm de SOITEC, avec une couche de SOI de 300nm, un BOX (pour « buried oxide ») de 800nm, et un substrat à haute résistivité. La fabrication commence avec les implantations pour réaliser les jonctions p - n et leurs zones de contact. Ensuite, les guides sont gravés, et les zones de contact siliciurées afin de réduire leurs résistances de contact. Enfin, les composants sont encapsulés par de l'oxyde et planarisés, avant de former les électrodes des modulateurs.

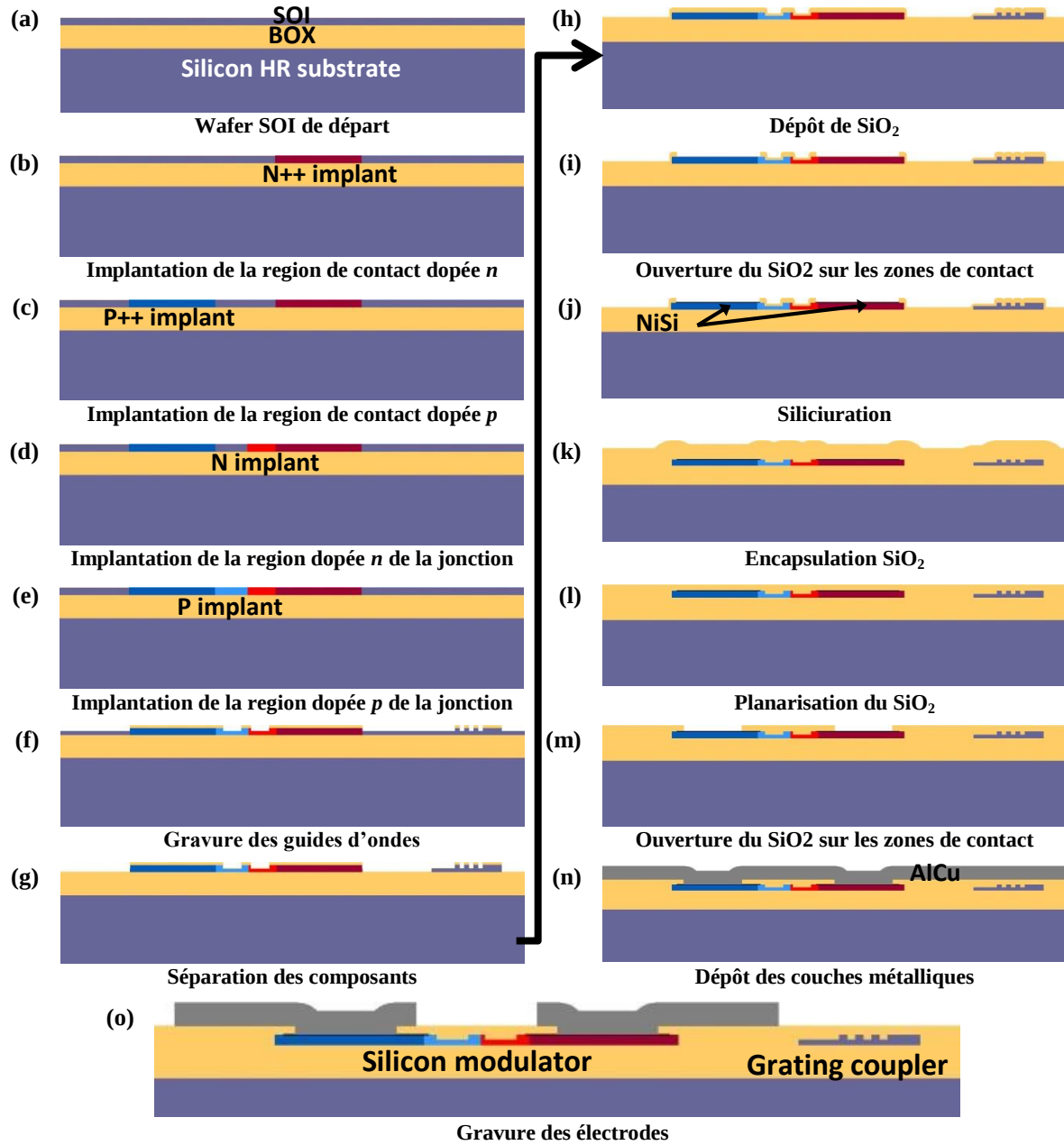


Figure A-22: Etapes de fabrication des modulateurs silicium.

Une fois leur fabrication terminée, les performances des différents composants (réseaux de couplage, virages, guides passifs, interféromètres multimodes, guides dopés) constituant les modulateurs sont caractérisées, à commencer par leurs performances optiques. Celles-ci sont résumées dans la **table A-1**, ainsi que celles des structures complètes de modulateurs. Bien qu'honorables, ces performances sont assez éloignées de celles de composants à l'état de l'art avec la même épaisseur de silicium (en particulier les réseaux de couplages et les guides passifs) [11]. Pour une prochaine démonstration, il serait nécessaire d'améliorer la conception des réseaux de couplages, ainsi que la rugosité des guides passifs. Néanmoins, il faut noter que les pertes dues au dopage de la jonction p - n sont estimées à 0.7dB/mm , ce qui est assez proche des 0.63dB/mm obtenus en simulation.

Table A-1: Pertes optiques mesurées pour chaque composant, et pour la structure complète de modulateur.

Composant	Pertes individuelles		Structure complète		Pertes totales pour chaque type de composant (dB)
	dB/par composant	dB/mm	Nombre de composants	Longueur (mm)	
Réseaux de couplage	4.2		2		8.4
Virages à 90°	0.06		14		0.84
Guides passifs		0.8		3.25	2.6
Interféromètres multimodes	0.1		2		0.2
Guides dopés		1.5		2 / 4 / 6	3 / 6 / 9
Pertes totales (somme des composants, pour une zone active de 2 / 4 / 6mm)					15,04 / 18,04 / 21,04
Pertes totales (moyenne des mesures, pour une zone active de 2 / 4 / 6mm)					15,4 / 18,4 / 21,4

L'étape suivante consiste à évaluer l'efficacité de la jonction $p-n$ implantée dans les modulateurs. Celle-ci est évaluée en mesurant le décalage en longueur d'onde des pics d'interférométrie du spectre des modulateurs pour différentes tensions (voir **figure A-23(a) à (c)**). En comparant ce décalage à l'écart entre deux pics (appelé intervalle spectral libre ou ISL), qui correspond à un déphasage de phase de 360° , il est possible d'évaluer le décalage de phase de la jonction en fonction de la tension appliquée, qui est comparé à celui obtenu par simulation (voir **figure A-23(d)**). De plus, la mesure de l'intervalle spectral libre permet aussi d'évaluer l'indice de groupe des guides d'onde, grâce à la formule $n_g = \lambda^2 / (ISL \times \Delta L)$, où ΔL correspond à la différence de longueur entre les deux bras du modulateur. Avec un ISL moyen de $4.2nm$, $\Delta L = 100\mu m$, et $\lambda = 1.3\mu m$, on obtient $n_g \approx 4$.

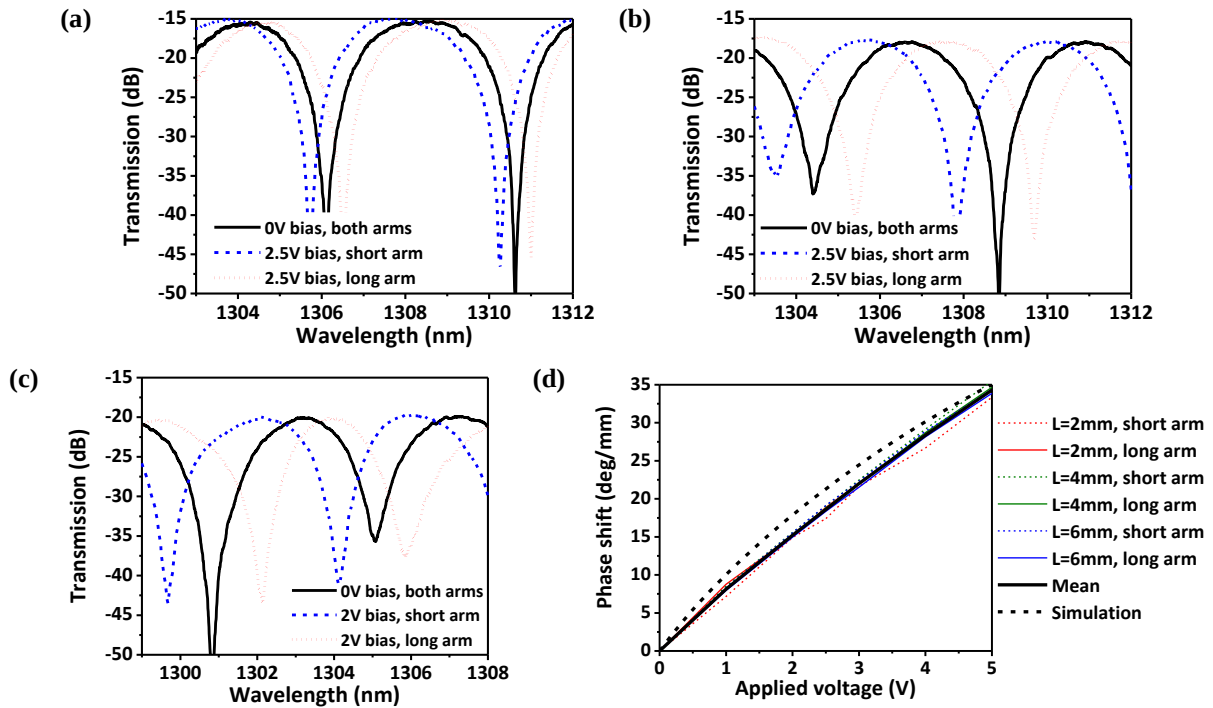


Figure A-23: Spectre de transmission des modulateurs à différentes tensions, depuis la fibre d'entrée, jusqu'à la fibre de sortie, pour des longueurs de zones actives de (a) 2mm, (b) 4mm, et (c) 6mm. (d) Décalage de phase en fonction de la tension inverse appliquée, comprenant les simulations et les mesures pour toutes les longueurs de zone active.

Les performances électro-optiques à hautes fréquences du modulateur sont aussi évaluées. Les réponses électro-optiques des modulateurs sont visibles sur la **figure A-24**. Pour une tension de polarisation inverse de 1.25V appliquée aux modulateurs, des bandes passantes de 28.2GHz, 23.1GHz, et 18GHz sont ainsi respectivement obtenues pour des longueurs de zones actives de 2mm, 4mm et 6mm. Bien que proches des simulations, ces différences s'expliquent par des pertes RF plus élevées sur les dispositifs réels, ainsi qu'un

indice effectif RF plus faible. Deux hypothèses permettent de justifier les différences entre les simulations et les mesures : 1) la résistivité réelle du substrat (qui est inconnue), et 2) la géométrie exacte des électrodes.

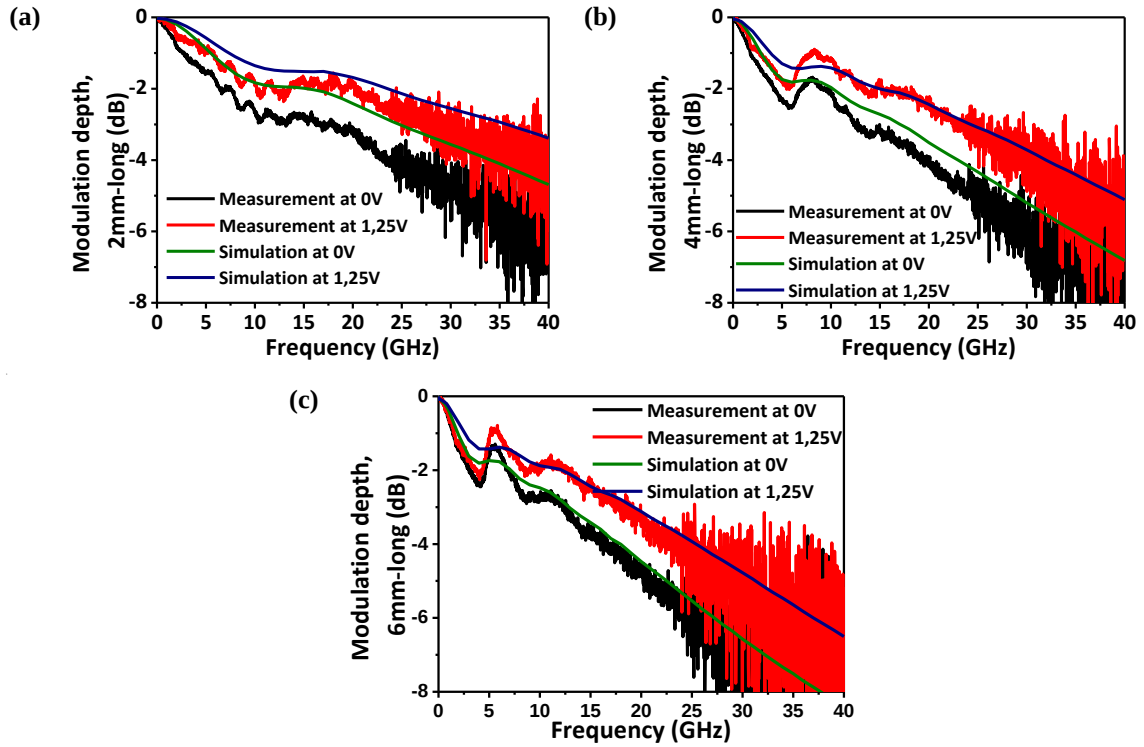


Figure A-24: Simulations et mesures des réponses électro-optiques pour différentes longueurs de zones actives: (a) 2mm, (b) 4mm, et (c) 6mm.

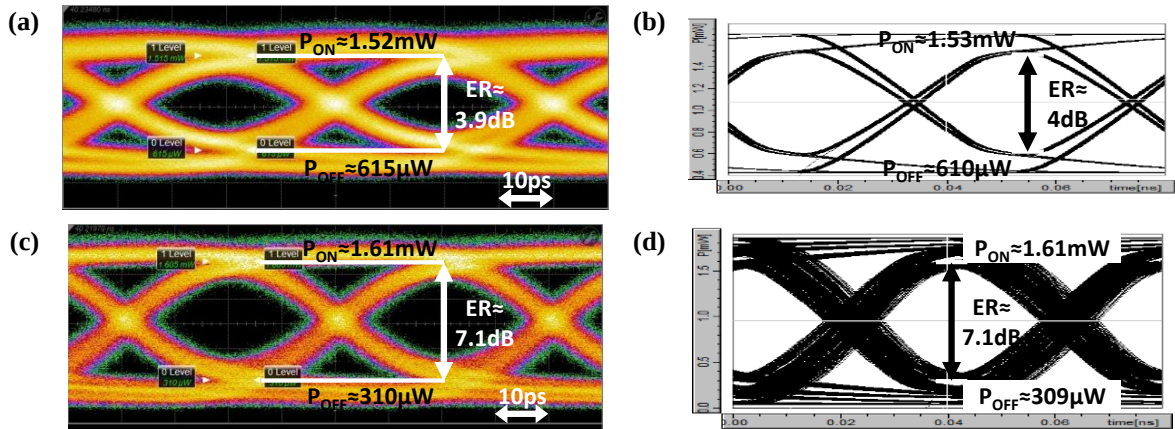


Figure A-25: Diagrammes de l'œil mesurés et simulés à 25Gb/s des modulateurs avec des zones actives de (a,b) 2mm et (c,d) 4mm, dans une configuration "dual drive push-pull", avec une tension d'amplitude 2.5V appliquée alternativement sur chaque bras, chacun polarisé à $-1.25V$.

La dernière phase de caractérisation consiste à étudier la réponse temporelle des modulateurs lorsqu'un signal électrique de 25Gb/s est utilisé comme signal de commande. Cette réponse est visualisée à l'aide de diagrammes de l'œil, où une séquence pseudo-aléatoire de bits "0" et "1" synchronisés et superposés les uns sur les autres. Puisque les bras des modulateurs peuvent être contrôlés séparément, il est préférable d'utiliser une configuration dite en « dual-drive push-pull », où un signal d'amplitude 2.5V est envoyé alternativement sur chaque bras du modulateur (polarisé à $-1.25V$), afin d'obtenir un taux d'extinction plus élevé, sans augmenter l'amplitude de la tension de commande des modulateurs. Ainsi, pour coder un bit de valeur "0", une tension de $-2.5V$ est envoyée sur l'un des bras pendant que l'autre est à 0V, et les tensions sur chaque bras sont inversés pour coder un bit de valeur "1". Dans cette configuration, il est préférable d'ajouter une différence de phase

statique de $\pi/2$ entre les deux bras du modulateur, afin de maximiser l'OMA du modulateur. Ce point particulier (nommé « quadrature ») offre ainsi le meilleur compromis entre taux d'extinction et pertes optiques du modulateur. Le modulateur étant asymétrique, il peut être placé en quadrature en choisissant la longueur d'onde du laser correspondant à une différence de phase statique de $\pi/2$.

Les diagrammes de l'œil électro-optiques à 25Gb/s des modulateurs de 2 et 4mm sont visibles sur la **figure A-25(a) et (c)**. Les deux diagrammes sont ouverts, avec des taux d'extinctions respectifs de 3.9dB et 7.1dB. Les paramètres statiques et RF des modulateurs obtenus lors des mesures précédentes sont aussi utilisés pour réaliser une simulation de diagramme de l'œil pour chacune des deux longueurs de zone active (voir **figure A-25(b) et (d)**). Le résultat de ces simulations correspond aux mesures effectuées, et permet d'évaluer plus précisément les performances des modulateurs qui sont ainsi résumées dans la **table A-2**. Bien que les performances statiques de nos modulateurs soient comparables à celles de l'état de l'art, ceux-ci souffrent de pertes élevées des composants passifs (réseaux de couplage, guides passifs). De plus, même si les bandes passantes électro-optiques de nos modulateurs sont moins élevées que celles de l'état de l'art, la conception des électrodes a été limitée par l'utilisation d'un unique niveau de métal, et pourrait encore largement être améliorée.

Table A-2: Résumé des performances des modulateurs de la fibre d'entrée jusqu'à la fibre de sortie.

Longueur de zone active (mm)	Pertes dues au dopage (dB/mm)	Décalage de phase à -2.5V (°/mm)	$V_{\pi}L_{\pi}$ à -2.5V (V.cm)	Maximum de transmission (dB)	Bande passante électro-optique à -1.25V (GHz)	Taux d'extinction à 25Gb/s avec 2.5V(dB)	Pertes d'insertion à 25Gb/s avec 2.5V (dB)
2	0.7	18.6	2.4	-15	28.2	3.9	16.4
4				-18	23.1	7.1	18.6

En utilisant les performances extraites de nos modulateurs, nous avons pu évaluer l'OMA pour différentes puissances optiques d'entrée (voir **figure A-27(a)**). Bien que le modulateur de 2mm soit supérieur en terme d'OMA, la différence avec le modulateur de 4mm reste faible. A cause des pertes optiques de nos composants, un laser devrait fournir des puissances élevées (supérieures à 17.5dBm ou 53.7mW) pour atteindre l'objectif fixé de -1.3dBm pour l'OMA, même sans tenir compte des pertes additionnelles dues aux 10km de transmission que l'on souhaite atteindre. Dans un cas où le laser serait intégré, les pertes de l'un des coupleurs seraient alors retirées (si l'on considère que l'intégration du laser ne perturbe pas les performances des autres composants), mais une puissance de 13.1dBm (ou 20.4mW) devrait encore être fournie par le laser dans le guide optique, ce qui reste très élevé. Afin de réduire la puissance laser nécessaire, il serait indispensable de réduire les pertes des composants passifs. En utilisant les pertes de composants passifs à l'état de l'art [11] (2.2dB pour les réseaux de couplage, 0.18dB/mm pour les guides passifs), tout en conservant la même conception de jonction *p-n* et d'électrodes pour le modulateur, une nouvelle évaluation de l'OMA a été réalisée (voir **figure A-27(b)**). Dans ce cas, la puissance requise pour atteindre l'objectif d'OMA de -1.3dBm est réduite à 9.7dBm (ou 9.3mW) pour une source laser externe, et même à 7.5dBm (ou 5.6mW) pour une source laser intégrée.

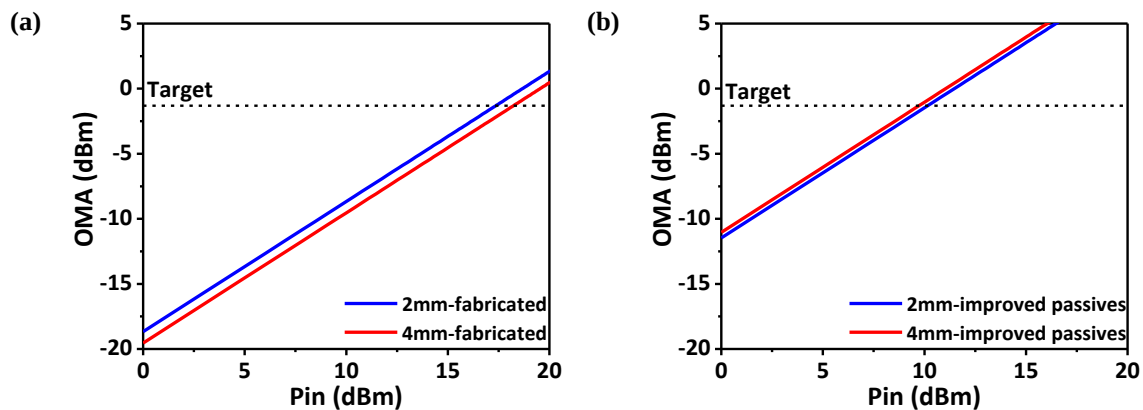


Figure A-27: OMA calculé à 25Gb/s pour (a) les modulateurs fabriqués, et (b) des modulateurs avec la même jonction *p-n* et les mêmes électrodes, mais de meilleurs composants passifs. Dans les deux cas, les pertes de propagation provenant d'une fibre monomode de 10km ne sont pas prises en compte.

A-3. Transmetteur hybride III-V sur silicium

Maintenant que la conception du modulateur silicium a été validée, l'étape suivante est d'étudier le laser hybride III-V sur silicium, afin de former le transmetteur intégré. La structure choisie pour le laser est celle d'un DBR, dont les réseaux de Bragg formant la cavité laser sont situés dans le guide silicium, de chaque côté du guide III-V. Cette structure permet d'obtenir un laser monomode, dont la longueur d'onde peut être contrôlée à l'aide d'une chaufferette située au-dessus des réseaux de Bragg.

Des schémas de la structure laser DBR utilisée sont visibles sur la **figure A-28**. L'empilement III-V est collé au-dessus d'un guide silicium, séparé par une membrane de 100nm de SiO₂. Cet empilement est basé sur le système InGaAsP/InP, qui est largement utilisé pour les transmissions longues distances. Le gain optique est fourni par des puits quantiques d'InGaAsP, dont le maximum de gain est centré sur 1.3μm. Les puits quantiques sont entourés par deux couches d'InGaAsP (appelée SCH pour « separate confinement heterostructure ») avec une indice de réfraction plus faible que celui des puits quantiques, afin de mieux confiner le mode optique dans la zone de gain. L'ensemble est lui-même entouré par deux couches d'InP dopées *n* et *p*, afin de former une jonction *p-i-n* pour injecter les porteurs de charges, et créer l'inversion de population nécessaire pour obtenir du gain optique. La couche dopée *n* est placée entre le III-V et le silicium, car celle-ci requiert un dopage moins élevé à sa surface pour obtenir un bon contact électrique avec le métal. A l'inverse, la zone dopée *p* requiert un dopage élevé à sa surface de contact, qui doit être éloignée de la zone de gain. Ainsi, la zone dopée *p* mesure généralement 1 à 2μm d'épaisseur, et l'empilement III-V total entre 2 et 3μm. Une fois gravé, le guide de III-V est long de 700μm et large de 5μm. Bien qu'il eut été intéressant d'étudier d'autres longueurs de zone active, l'empreinte totale du transmetteur étant trop grande, une seule longueur fut étudiée. De même, un guide moins large aurait pu permettre de réduire le courant de seuil du laser, mais celui-ci serait aussi devenu plus sensible aux pertes de dispersion causées par la rugosité du guide III-V après gravure, et son contact métallique aurait été plus difficile à aligner sur le guide.

Afin de transférer la lumière du guide III-V vers le guide silicium en utilisant un couplage adiabatique, il est nécessaire de réaliser une inversion du ratio des indices effectif des deux guides sur la zone de couplage. Ainsi, l'indice effectif du guide III-V doit être le plus élevé dans la zone centrale du guide III-V afin de maximiser le confinement optique du supermode pair dans la zone de gain, et l'indice effectif du guide silicium doit être le plus élevé aux bords du guide III-V afin de maximiser le confinement optique du supermode pair dans le guide silicium. Cette variation d'indice est réalisée à l'aide d'un taper situé dans le guide silicium à chaque extrémité du guide III-V (voir **figure A-28(c)**). Cependant, étant donné la géométrie du guide III-V, son indice effectif est trop élevé (≈ 3.29) pour réaliser une inversion du ratio d'indice avec un guide silicium de 300nm d'épaisseur comme celui utilisé pour le modulateur silicium. Par conséquent, il est nécessaire d'utiliser un guide silicium plus épais, qui est fixé à 500nm dans notre cas. Ce guide est ensuite gravé de 200nm, laissant ainsi un « slab » de 300nm, utilisé pour former les autres composants du transmetteur.

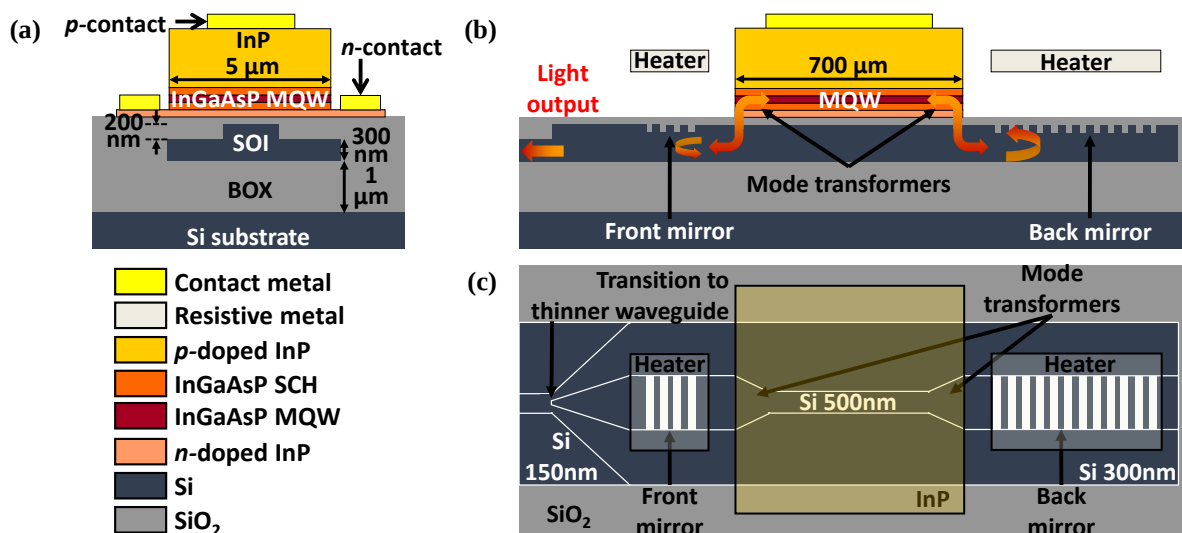


Figure A-28: Vues schématiques (a) transverse, (b) longitudinale, et (c) de dessus du laser.

Afin de transférer efficacement la lumière entre les deux guides sur une distance relativement courte (fixée ici à $100\mu\text{m}$), les tapers sont conçus à l'aide du critère d'adiabaticité proposé par Sun et al. [79], [80], qui permet de limiter la fraction de puissance dispersée dans les autres modes optiques (ε) pendant le transfert:

$$\gamma = \frac{\delta(z)}{\kappa} = \tan \left[\arcsin \left(2\kappa \varepsilon^{1/2} [z - z_0] \right) \right] \quad (\text{A.2})$$

Où γ est un paramètre de normalisation, δ est la différence de constantes de propagation entre les modes des guides III-V et silicium, κ est la constante de couplage entre les deux guides, z est la position du mode optique, et z_0 correspond à la position où les guides sont en accord de phase ($\gamma = 0$). D'après cette formule, si $(z - z_0)$ varie de $-1/(2\kappa \varepsilon^{1/2})$ à $1/(2\kappa \varepsilon^{1/2})$, alors γ varie de $-\infty$ à $+\infty$, signifiant que la puissance du supermode pair est effectivement transférée entre les deux guides. Ainsi, un taper respectant ce critère, avec une longueur maximum $L_{ac} = 1/(\kappa \varepsilon^{1/2}) = 2z_0$, permet un transfert du supermode pair avec une fraction de puissance ε dispersée dans le supermode impair.

Afin de réaliser ce taper, il faut exprimer la largeur du guide W_{Si} en fonction du paramètre γ , puis utiliser le critère d'adiabaticité pour trouver la forme du taper. Le paramètre $\delta(W_{\text{Si}})$ est extrait des indices effectifs des modes fondamentaux TE (pour « transverse electric ») des deux guides isolés (voir **figure A-29**), obtenus par des simulations à éléments finis. Dans notre cas, l'accord de phase est réalisé lorsque la largeur du guide silicium est de $0.85\mu\text{m}$.

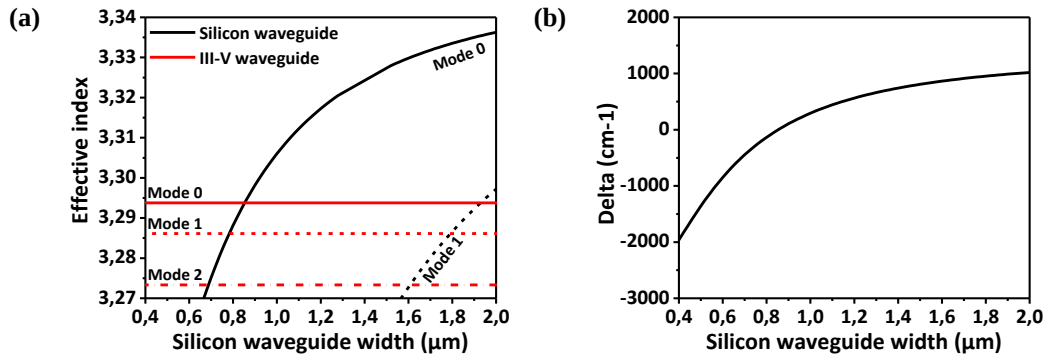


Figure A-29: (a) Indices effectifs des modes TE des guides III-V silicium en fonction de la largeur de l'arête du guide silicium, à $1.31\mu\text{m}$. Les deux guides sont encapsulés dans de l'oxyde. (b) Différence de constantes de propagation entre les modes fondamentaux de chaque guide.

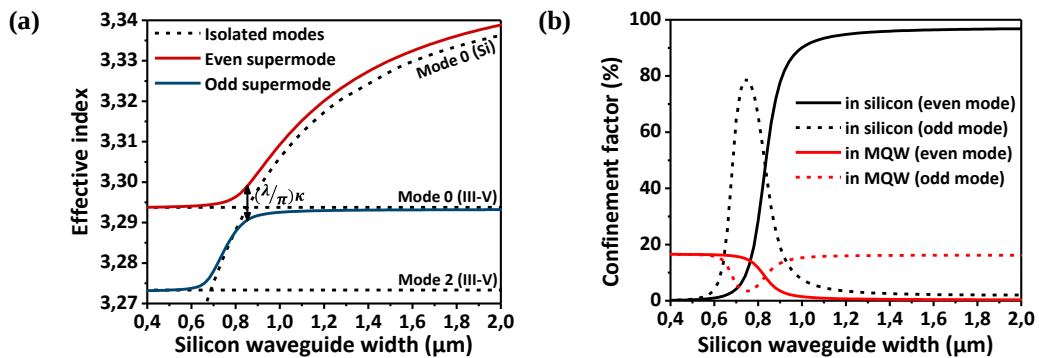


Figure A-30: (a) Indices effectifs des supermodes pairs et impairs du guide hybride III-V sur silicium, et (b) facteur de confinement de chaque supermode dans le guide de silicium, et dans les puits quantiques (sans les barrières séparant les puits), en fonction de la largeur du guide silicium à $1.31\mu\text{m}$.

La constante de couplage κ dépend principalement de la distance séparant les deux guides, et peut-être extrait des indices effectifs des supermodes pairs et impairs (voir **figure A-30(a)**). En effet, celui-ci est extrait à la largeur de guide correspondant à l'accord de phase, et vaut $\kappa = (\pi/\lambda)(\tilde{n}_e - \tilde{n}_o) = 201.45\text{cm}^{-1}$. Bien que les supermodes suivent les modes TE des guides isolés, ils ne se superposent pas parfaitement après l'accord de phase, ce qui signifie que le couplage est relativement fort, alors que la théorie de couplage est normalement

limitée à des couplages faibles. Par conséquent, la valeur de κ est en réalité plus élevée, mais cela ne pose pas de problèmes pour la conception du taper. Le facteur de confinement des supermodes dans le guide silicium et dans les puits quantiques a aussi été évalué (voir **figure A-30(b)**). Le maximum de confinement du mode pair dans les puits quantiques est estimé à 16.5%. Quand la largeur du guide silicium augmente, le mode pair passe du guide III-V au guide silicium, alors que le mode impair reste principalement dans le guide III-V. Comme les réflecteurs sont situés dans le guide silicium, le mode impair aura besoin de beaucoup plus de courant pour atteindre le seuil laser. La majorité du transfert du mode pair se produit pour des largeurs de guide silicium comprises entre $0.6\mu\text{m}$ et $1.2\mu\text{m}$.

Une fois les dépendances des paramètres δ et κ sur la largeur du guide silicium sont connues, il est possible de déterminer la dépendance $W_{\text{Si}}(\gamma)$ à l'aide d'une fonction de fit, et d'utiliser le critère d'adiabaticité pour trouver la dépendance $W_{\text{Si}}(z)$ à une valeur ε fixée. De plus, puisque la valeur réelle de κ est plus élevée que sa valeur théorique, la valeur réelle de ε sera en fait plus faible que sa valeur théorique. Ensuite, il suffit de choisir un couple de largeur d'entrée et de sortie du taper, afin de fixer sa forme. Un exemple de forme normalisée de taper (utilisé pour le laser DBR) est visible sur la **figure A-31(a)**. La dernière étape consiste à affiner les largeurs d'entrée et de sortie du taper pour une longueur fixée de $100\mu\text{m}$. Comme on peut le voir sur la **figure A-31(b)**, plusieurs couples de largeurs permettent d'obtenir une efficacité supérieure à 90%. Finalement, la largeur d'entrée du taper est fixée à $0.6\mu\text{m}$, et sa largeur de sortie à $1.55\mu\text{m}$. Bien que ce taper soit robuste vis-à-vis des variations de largeurs, il reste sensible aux variations de l'épaisseur de la couche de SiO_2 entre les deux guides, une variation de $\pm 30\text{nm}$ suffisant à faire chuter l'efficacité de couplage sous 70%.

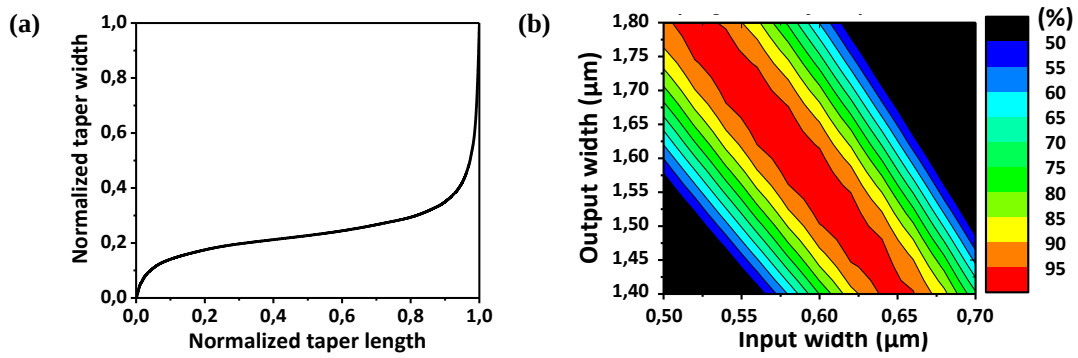


Figure A-31: (a) Forme normalisée d'un taper adiabatique. (b) Facteur de confinement de la puissance optique dans le guide de silicium à l'extrémité du guide III-V en fonction des largeurs d'entrée et de sortie du taper adiabatique.

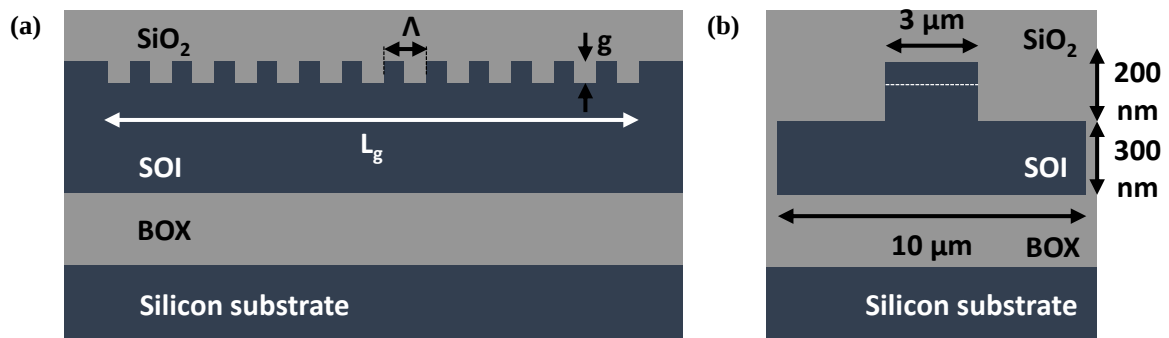


Figure A-32: Vues schématiques (a) longitudinales and (b) transversales des miroirs de Bragg.

Une fois la lumière transférée dans le guide silicium, celle-ci est réfléchiée par deux réseaux de Bragg situés dans le guide silicium de chaque côté du guide III-V, formant ainsi la cavité laser illustrée sur la **figure A-28(c)**. Deux tapers linéaires de $60\mu\text{m}$ de long sont utilisés pour connecter la sortie des tapers adiabatiques à des guides de $3\mu\text{m}$ de large où sont définis les réseaux de Bragg. Ces réseaux doivent non seulement agir en tant que réflecteurs, mais aussi en tant que filtre en longueur d'onde afin d'obtenir un laser monomode. Afin d'obtenir un bon confinement dans la zone de gain mais aussi d'extraire la lumière dans une seule direction, le miroir arrière doit être aussi réfléchissant que possible, alors que la réflectance du miroir avant doit être plus faible. Les miroirs

de Bragg sont représentés sur la **figure A-32**. La largeur des miroirs étant fixée à $3\mu\text{m}$ et celle du rapport cyclique des réseaux à 0.5, les performances des réseaux sont alors contrôlés par trois paramètres géométriques : la profondeur de gravure des réseaux (g), leur longueur (L_g) et leur période (Λ).

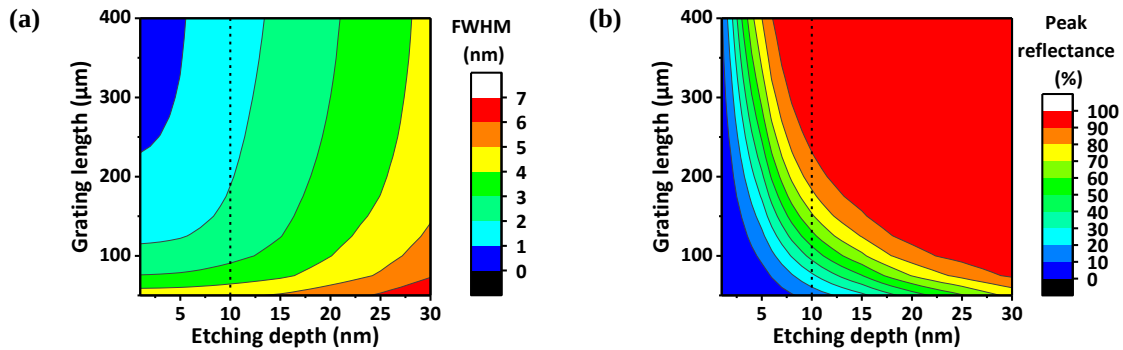


Figure A-33: (a) FWHM et (b) réflectance maximum en fonction de la profondeur de gravure et la longueur des réseaux de Bragg. Les réseaux sont centrés sur la longueur d'onde de 1310nm .

La profondeur de gravure influence principalement la sélectivité en longueur d'onde des miroirs de Bragg, ainsi que leur réflectance maximum. Plus cette profondeur est élevée, plus la différence d'indice effectif entre les parties gravées et non-gravées est importante, et plus la réflectance du laser augmente. Cependant, le miroir perd alors en sélectivité en longueur d'onde. Cette sélectivité se mesure à l'aide de la largeur à mi-hauteur (« full width at half maximum » en anglais ou encore FWHM) du filtre en longueur d'onde. Afin d'obtenir un laser monomode, il est donc nécessaire d'utiliser une faible profondeur de gravure, avec le risque de devoir utiliser des réseaux extrêmement longs pour obtenir une réflectance proche de 100%. L'influence de la profondeur de gravure, ainsi que de la longueur du réseau sur sa FWHM et sur sa réflectance maximum sont visibles sur la **figure A-33**. Finalement, un compromis est trouvé pour la profondeur de gravure est fixée à 10nm . La longueur du réseau influence principalement la réflectance des miroirs. Cependant, à longueur de miroir fixe, une variation de 1nm sur la profondeur de gravure peut avoir une large influence sur la réflectance. Afin d'obtenir un miroir arrière avec une réflectance proche de 100%, sa longueur est fixée à $700\mu\text{m}$, alors que celle du miroir avant est fixée à $150\mu\text{m}$ pour viser une réflectance maximum comprise entre 35% et 85% pour une variation de $\pm 4\text{nm}$ sur la profondeur de gravure.

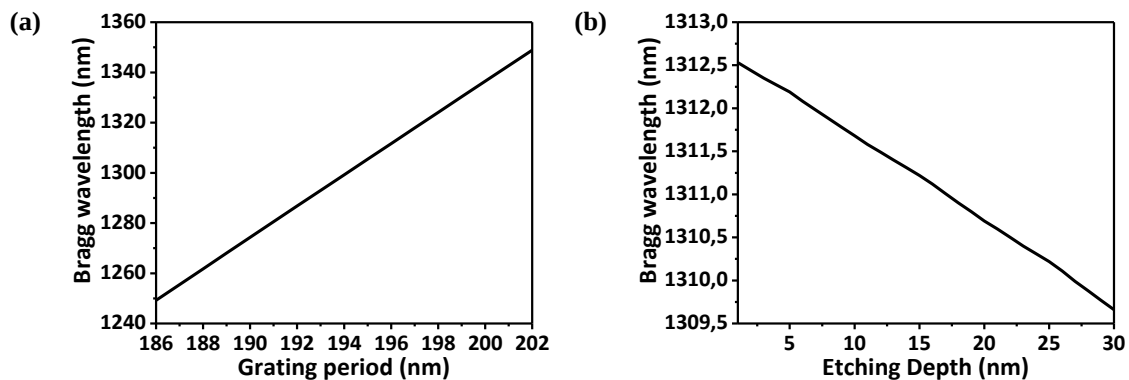


Figure A-34: Longueur d'onde de Bragg en fonction de (a) la période du réseau de Bragg pour une profondeur de gravure de 10nm , et (b) de la profondeur de gravure pour une période de 196nm .

Le dernier paramètre, la période du réseau de Bragg n'a pas d'influence sur sa sélectivité ou sa réflectance, mais permet de contrôler la longueur d'onde sur laquelle il est centré (nommé longueur d'onde de Bragg ou λ_g). Les deux sont liés par la relation $\lambda_g = 2\bar{n}\Lambda$, où $\bar{n}$ est la valeur moyenne de l'indice effectif des zones gravées et non-gravées. La dépendance de la longueur d'onde de Bragg sur la période pour une profondeur de gravure de 10nm est ainsi illustrée sur la **figure A-34(a)**. Ainsi, une variation de $\pm 1\text{nm}$ sur la période du réseau, résulte en une variation de $\pm 6.2\text{nm}$ sur la longueur d'onde de Bragg. Afin de respecter l'objectif fixé de moduler quatre longueurs d'ondes centrées autour de $1.3\mu\text{m}$ avec un espacement de 4.5nm , il sera donc nécessaire d'utiliser une

chaufferette pour modifier l'indice moyen du réseau et adapter ainsi la longueur d'onde de Bragg. La profondeur de gravure a aussi une influence sur la longueur d'onde de Bragg (voir **figure A-35(b)**), mais qui est assez faible, puisqu'une variation de $\pm 1nm$ sur la profondeur de gravure résulte en une variation de $\pm 0.1nm$ sur la longueur d'onde de Bragg. Finalement, les caractéristiques des réseaux choisis sont résumées dans la **table A-3**.

Table A-3: Caractéristiques des réseaux de Bragg.

Désignation	Longueur (μm)	Profondeur de gravure (nm)	Réflectance maximum (%)	FWHM (nm)	Période (nm)	Longueur d'onde de Bragg (nm)
Miroir avant	150	10	≈ 70	2.2	195	1303.3
					197	1315.7
Miroir arrière	700		≈ 100	1.4	195	1303.3
					197	1315.7

Une vue globale du transmetteur est visible sur la **figure A-35**. Une fois sortie de la cavité laser, la lumière est ensuite transférée à l'aide d'un taper exponentiel d'un guide de 500nm d'épaisseur (avec un slab de 300nm), vers un guide de 300nm d'épaisseur (avec un slab de 150nm), où sont définis les autres composants formant le transmetteur. La lumière traverse ensuite le modulateur silicium, avant d'être couplé dans une fibre monomode à l'aide d'un réseau de couplage. Des lasers avec des périodes de réseaux de 195nm et 197nm ont été associés à des modulateurs avec des zones actives de 2mm, 4mm et 6mm, mais seulement ceux de 2mm et 4mm ont été testés. Deux paires de chaufferettes sont aussi situées au-dessus des réseaux de Bragg du laser et de la zone passive du modulateur. La première paire permet ainsi de contrôler la longueur d'onde d'émission du laser, alors que la seconde permet d'ajuster la différence statique entre les bras du modulateur pour le mettre en quadrature pendant les mesures.

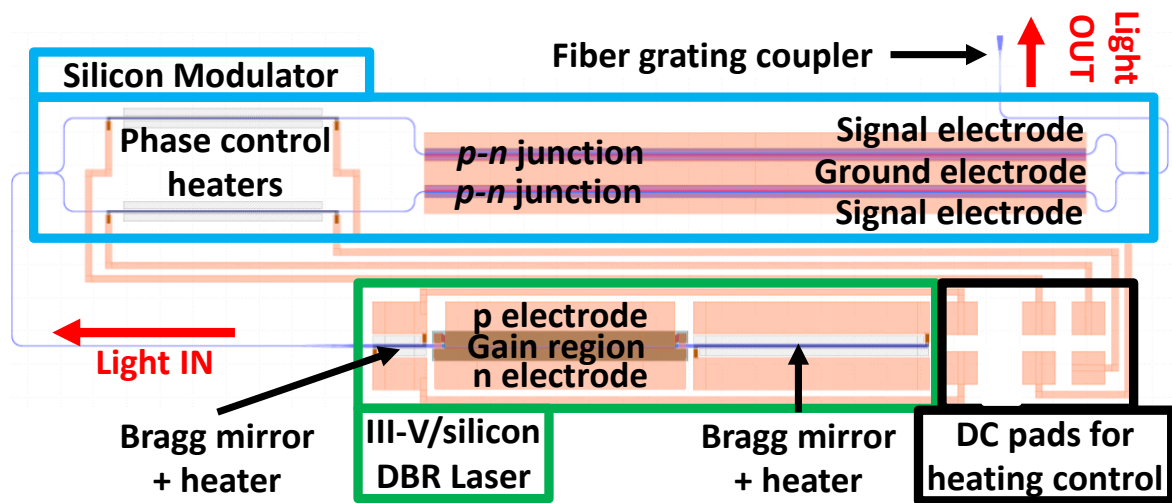


Figure A-35: Masques du transmetteur complet, intégrant à la fois un laser hybride III-V sur silicium et un modulateur silicium avec une zone active longue de 2mm.

Les transmetteurs ont été fabriqués en utilisant les étapes de fabrications décrites sur la **figure A-36**. La fabrication débute sur des substrats SOI de 200mm de SOITEC, avec une couche de SOI de 500nm, un BOX de 1 μm , et un substrat à haute résistivité. Premièrement, les réseaux de Bragg sont gravés, suivis par les guides et tapers adiabatiques utilisés pour les lasers. Ensuite les étapes décrites sur la **figure A-22** sont réutilisées pour fabriquer les modulateurs et les composants passifs, jusqu'à la planarisation des composants avant collage, avec une épaisseur visée de 100nm au-dessus des guides lasers. Cependant, pour des raisons de non-uniformité des dépôts et de la planarisation, cette épaisseur n'a été respectée qu'au centre des plaques, alors que l'épaisseur diminue jusqu'à 20nm aux bords. Comme une seule plaque de SOI était disponible, plusieurs plaques de III-V de 50mm ont été collées sur cette même plaque, mais pas en son centre, où l'épaisseur de collage était optimisée. La composition de ces plaques (fournies par Landmark) est détaillée sur la **table A-4**.

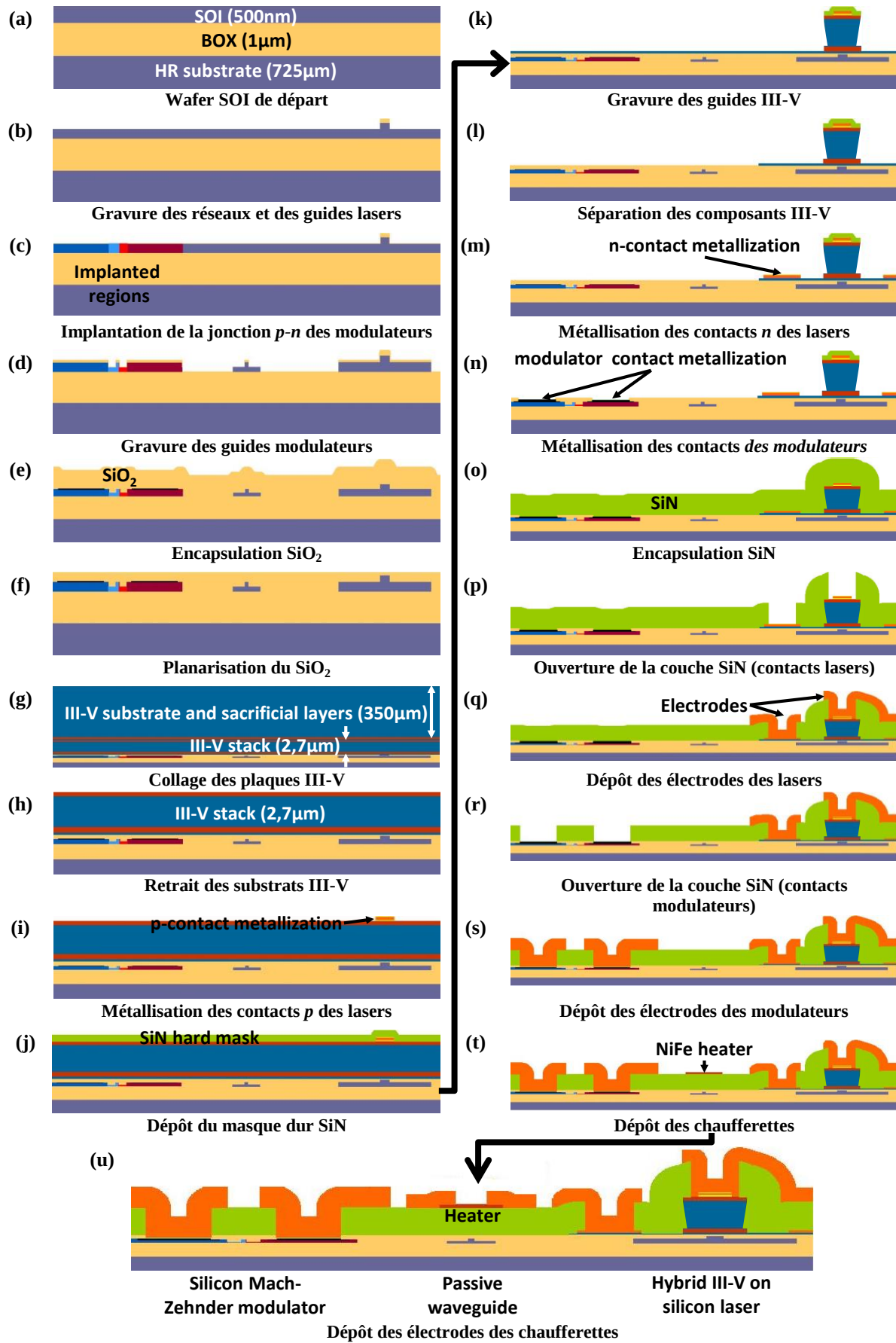


Figure A-36: Etapes de fabrication des transmetteurs hybrides III-V sur silicium.

Après le collage, le substrat, ainsi que les différentes couches d'arrêt et sacrificielles sont retirées chimiquement. Les plaques de SOI sont ensuite réduites de 200mm à 75mm autour des plaque de III-V, afin de poursuivre la fabrication dans une autre salle blanche dédiée aux procédés de fabrication III-V. Les contacts métalliques p des lasers sont formés par lift-off (tous les contacts et les électrodes sont d'ailleurs formés par cette méthode), suivis des guides III-V qui sont définis en alternant des gravures sèches (pour les couches d'InGaAs et InGaAsP) et humides (pour les couches d'InP). Après avoir formé les contacts n des lasers, l'oxyde au-dessus des zones de contact des modulateurs est ouvert, afin d'y déposer une couche de contact métallique. Une épaisse couche de SiN ($2\mu\text{m}$) est ensuite déposée pour encapsuler les structures. Cette couche est ensuite ouverte par gravure sèche au niveau des contacts du laser, puis les électrodes sont formées avec une couche d'or épaisse d' $1.3\mu\text{m}$ et la méthode du lift-off. Les mêmes étapes sont répétées pour former les électrodes du modulateur. Pour finir, les chauffetterettes sont formée par le lift-off d'une couche de 250nm de NiFe, connectée par des électrodes d'or épaisses de 250nm. Un transmetteur complet est illustré sur la **figure A-37**, ainsi que les différents éléments formant ce transmetteur, et présentés précédemment.

Table A-4: Composition des plaques de III-V utilisées pour les transmetteurs hybrides III-V sur silicium. Les couches grisées sont retirées après le collage, et seules les autres couches sont utilisées pour former le laser.

Couche	Matériau	Pic de photoluminescence (μm)	Bandgap (eV)	Épaisseur (nm)	Dopage (cm^{-3})
Substrat	InP			350 (μm)	2×10^{18}
Transition	InP			50	Non dopée
Couche d'arrêt	InGaAs			300	Non dopée
Couche sacrificielle	InP			300	Non dopée
Contact dopé p	InGaAs	1.65	0.75	200	2×10^{19}
Transition	InGaAsP	1.1	1.13	50	5×10^{18}
Couche dopée p	InP	0.92	1.35	2000	$2 \times 10^{18} \rightarrow 5 \times 10^{17}$
SCH	InGaAsP	1.1	1.13	100	Non dopée
Barrières (x7)	InGaAsP	1.1	1.13	10	Non dopée
Puits (x8)	InGaAsP	1.29	0.96	8	Non dopée
SCH	InGaAsP	1.1	1.13	100	Non dopée
Contact dopé n	InP	0.92	1.35	110	3×10^{18}
Super-lattice (x2)	InGaAsP	1.1	1.13	7.5	3×10^{18}
Super-lattice (x2)	InP	0.92	1.35	7.5	3×10^{18}
Interface de collage	InP	0.92	1.35	10	Non dopée

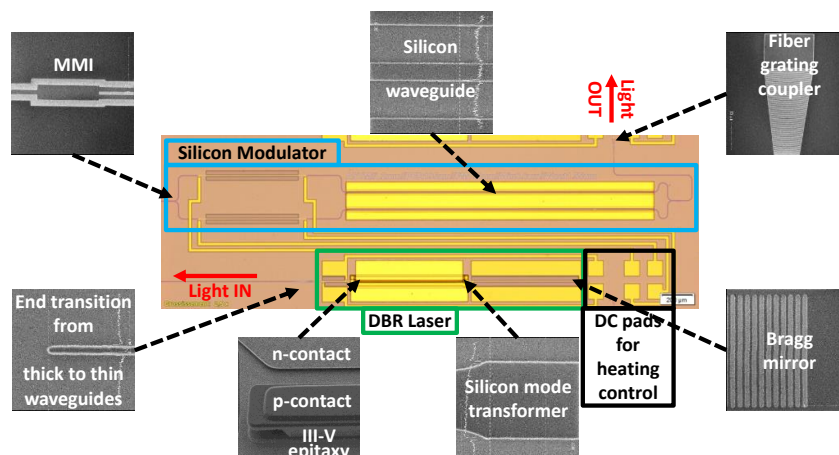


Figure A-37: Vue de dessus d'un transmetteur complet, avec des images prises au microscope électronique à balayage des différents éléments composant le transmetteur.

Une fois la fabrication des transmetteurs terminée, les différents composants les constituants sont d'abord caractérisés. Cependant, différents problèmes de fabrication ont été mis en évidence parmi ces mesures, et une évaluation précise des performances de tous les composants (comme celle réalisée pour les modulateurs) s'est

révélée impossible. Ainsi, les pertes optiques de tous les composants passifs sont supérieures à celles mesurées dans le cas du modulateur seul. La cause de cette augmentation pourrait provenir des différences d'épaisseur de SOI utilisés dans les deux cas, puisque dans le cas des transmetteurs, les guides de $300nm$ sont formés à partir d'une de la gravure d'une couche de $500nm$. Cette gravure supplémentaire pourrait avoir augmenté suffisamment la rugosité de surface des guides, et ainsi augmenter les pertes optiques. Les tests de structures laser sans modulateurs montrent l'apparition d'un courant de seuil laser $\approx 50mA$, mais avec des puissances émises dans une fibre monomode inférieures à $1mW$ (estimées au mieux à $1-3mW$ dans le guide silicium). Ces puissances réduites pourraient s'expliquer par l'épaisseur de collage entre les plaques SOI et III-V, estimée à $\approx 60nm$, au lieu des $100nm$ prévus dans la conception des lasers, réduisant ainsi leur efficacité de couplage dans le silicium. L'étude des spectres des lasers montrent que ceux-ci sont monomodes (avec un SMSR supérieur à $30dB$) jusqu'à des courants de $100mA$, puis deviennent multimodes ensuite. Cet effet provient de la conception des lasers, et pourrait être résolu en réduisant la longueur de la zone active III-V. Selon leur période de réseau, les lasers émettent à des longueurs d'onde proches de $1303.4nm$ (période de $195nm$), ou de $1315nm$ (période de $197nm$), ce qui est proche des longueurs d'ondes attendues (voir **table A-3**). Quant aux structures de modulateurs, bien que leur décalage de phase en tension soit resté proche de celui obtenu précédemment, leurs réponses en fréquence se sont dégradées, vraisemblablement à cause d'un problème de contact dû à une couche d'oxyde native entre les régions siliciurées et leur couche de contact métallique.

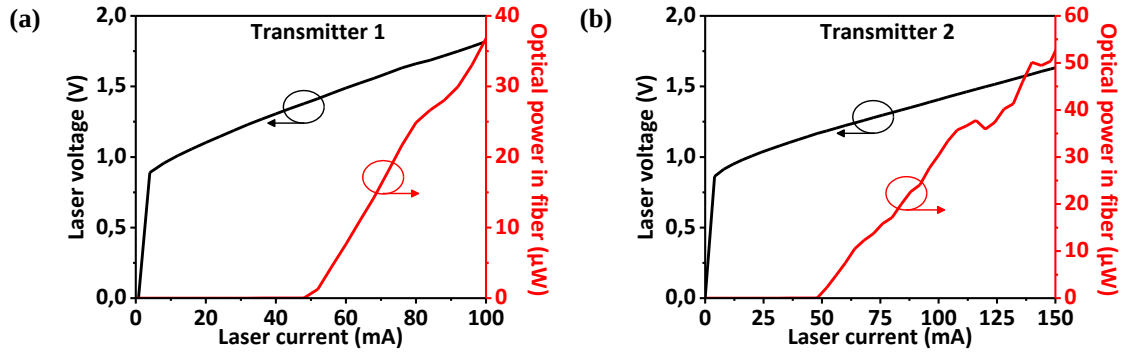


Figure A-38: Tension aux bornes du laser et puissance optique collectée dans la fibre en fonction du courant injecté dans le laser pour (a) le transmetteur 1 et (b) le transmetteur 2. Pour les valeurs de courant affichées, les lasers sont monomodes.

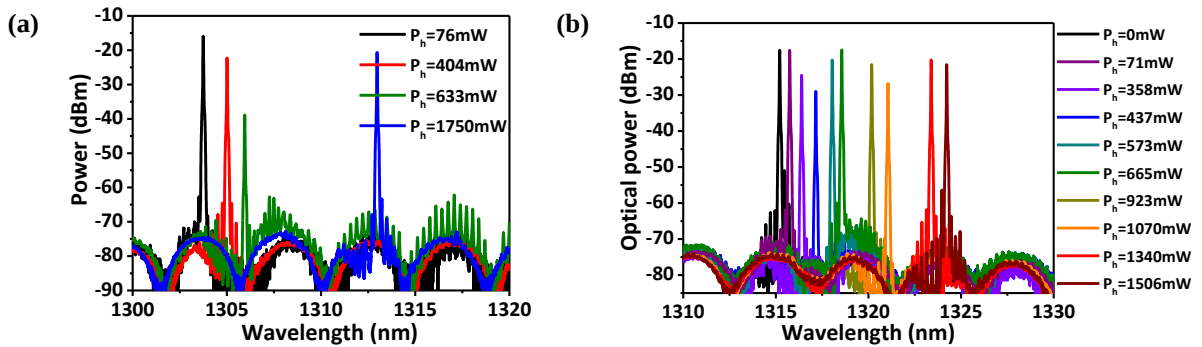


Figure A-39: Spectres pour une différence de phase fixe entre les bras des modulateurs, avec une longueur d'onde laser contrôlée (a) pour le transmetteur 1, avec un courant laser de $100mA$, et (b) pour le transmetteur 2, avec un courant laser de $150mA$. P_h est la puissance totale injectée dans les chaufferettes.

Deux transmetteurs complets ont finalement été testés : l'un avec une période de réseau de Bragg de $195nm$ et un modulateur de $4mm$ de long (transmetteur 1), et l'autre avec une période de réseau de Bragg de $197nm$ et un modulateur de $2mm$ de long (transmetteur 2). Les performances statiques de ces transmetteurs sont visibles sur la **figure A-38** (la différence de phase entre les bras des modulateur est réduite à 0 à l'aide des chaufferettes). Les courants de seuil des lasers sont à $48mA$, avec des puissances collectées dans la fibre respectives de 37 et $52\mu W$ pour les transmetteurs 1 et 2. Les spectres de ces transmetteurs sont aussi étudiés et illustrés sur la **figure A-39**. Les lasers sont monomodes, et leur émission spontanée permet de visualiser le spectre interférométrique

du modulateur. En augmentant la puissance injectée dans les chaufferettes situées au-dessus des réseaux, la longueur d'onde augmente de façon non-continue, et saute d'un mode à un autre. La longueur d'onde d'émission du laser peut ainsi être augmentée jusqu'à 8.5nm . Cependant, à cause des $2\mu\text{m}$ de SiN séparant les chaufferettes des réseaux, une puissance électrique élevée est requise dans les chaufferettes.

Les réponses temporelles des transmetteurs à 25Gb/s sont visibles sur la figure A-40. Les modulateurs sont en configuration « dual-drive push-pull », avec un signal d'amplitude 2.5V envoyé alternativement sur chaque bras du modulateur (polarisé à -1.25V). Les modulateurs sont aussi placés en quadrature à l'aide des chaufferettes situées au-dessus de leurs zones passives. Pour les transmetteur 1 et 2, le courant du laser est fixé respectivement à 100mA et 150mA , avec une longueur d'onde d'émission respective de 1303.5nm et de 1315.8nm . Dans les deux cas, les diagrammes de l'œil sont ouverts, avec des taux d'extinction de 4.7dB (transmetteur 1) et 2.9dB (transmetteur 2). Les différences avec les diagrammes de l'œil précédents s'expliquent à la fois par la réduction de bande passante des modulateurs et par l'utilisation d'un photodétecteur non présent dans la mesure précédente, avec une sensibilité plus élevée, mais une bande passante de 12GHz . Pour finir, une mesure similaire est réalisée avec une fibre monomode de 10km , dont les pertes additionnelles sont compensées par un amplificateur optique (voir figure A-41). Les diagrammes de l'œil obtenus sont presque identiques, avec les mêmes taux d'extinction, sans signe apparent de distorsion du signal, ce qui prouve que la transmission à 10km n'est limitée que par la puissance de sortie du transmetteur.

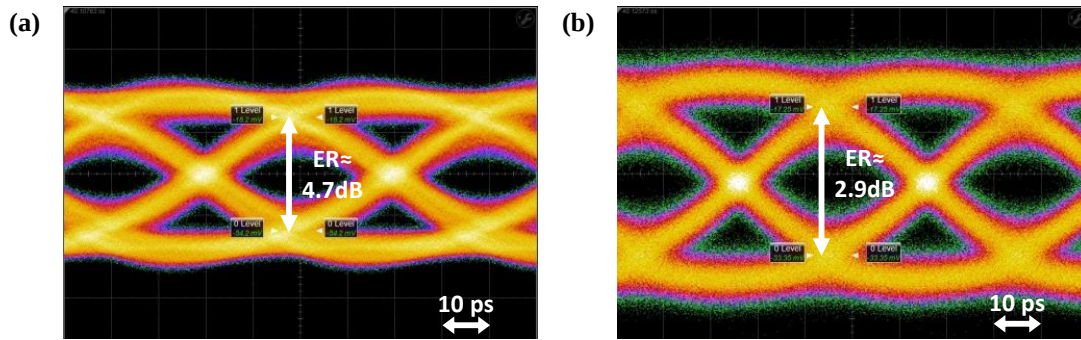


Figure A-40: Diagrammes de l'œil à 25Gb/s des transmetteurs hybrides III-V sur silicium avec (a) un modulateur de 4mm à 1303.5nm (transmetteur 1) et (b) un modulateur de 2mm à 1315.8nm (transmetteur 2), en « dual drive push-pull », avec une tension d'amplitude 2.5V sur chaque bras.

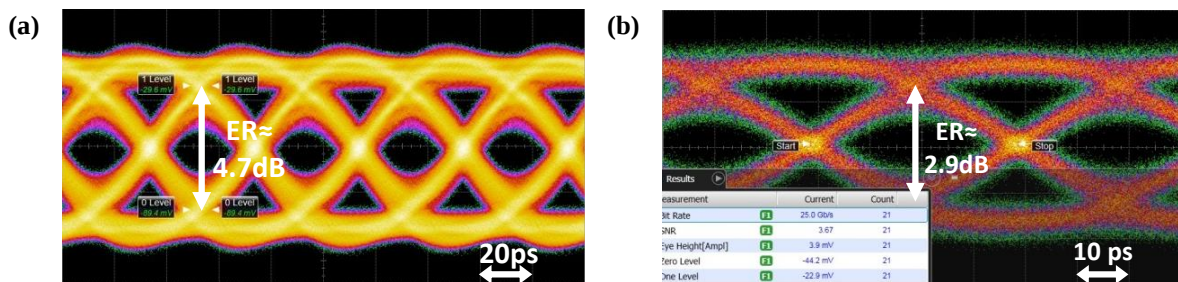


Figure A-41: Diagrammes de l'œil à 25Gb/s du (a) transmetteur 1 et (b) transmetteur 2, dans les mêmes conditions que la figure A-40, mais avec une fibre monomode de 10km de long.

Ces deux transmetteurs avec une source laser intégrée pour la photonique sur silicium sont donc capables de transmettre des informations à 25Gb/s sur des distances de 10km . Ces transmetteurs fonctionnent à deux longueurs d'ondes différentes autour de $1.3\mu\text{m}$, et d'autres longueurs d'ondes sont accessibles en changeant à la fois la période des réseaux de Bragg (pour un contrôle grossier de la longueur d'onde) et en utilisant des chaufferettes au-dessus de ces réseaux (pour un contrôle fin). Si l'on compare ces résultats aux objectifs visés à la fin de la section A-1, on peut voir que même si la plupart de ces objectifs ont été remplis, la puissance optique de sortie des transmetteurs reste largement insuffisante (l'OMA des deux transmetteurs étant estimée inférieure à -19dBm avec une fibre de 10km). Il est donc essentiel d'augmenter la puissance des lasers, mais aussi de réduire les pertes optiques des différents composants, comme l'a montré l'étude sur l'OMA à la fin de la section A-2.

A-4. Laser hybride III-V sur silicium amorphe

Bien que des transmetteurs avec une source laser intégrée sont démontrés dans la section précédente, ceux-ci sont basés sur une couche SOI plus épaisse (500nm) pour coupler la lumière des lasers dans le reste du circuit photonique. Cette couche est ensuite gravée de 200nm pour former les autres composants. Bien que suffisant pour une démonstration, cette approche n'est pas adaptée pour les plateformes de fabrication basées sur une fine couche de SOI (généralement inférieure à 300nm), puisqu'elle modifie la fabrication des composants déjà développés. De plus, cette gravure supplémentaire risque d'induire une rugosité supplémentaire à la surface des guides, ainsi qu'une variation d'épaisseur risquant d'influencer les performances des composants (qui pourrait expliquer en partie les pertes plus importantes des composants co-intégrés avec un laser). Plutôt que de définir une nouvelle plateforme de fabrication et de concevoir de nouveaux composants basés sur une épaisseur de SOI plus élevée, il est plus intéressant de trouver une solution pour coupler la lumière des lasers hybrides III-V sur silicium dans des guides fabriqués à partir d'une fine couche de SOI. Pour des raisons de compatibilité, les étapes liées au laser devraient être intégrées à la fin de la fabrication, et seront donc limitées à des budgets thermiques plus faibles (limités à $400\text{-}450^\circ\text{C}$ pendant les étapes de métallisation).

Pour résoudre ce problème, plusieurs solutions basées sur la réduction de l'indice effectif du guide III-V ont été proposées, afin d'atteindre un accord de phase avec un guide silicium plus fin. La solution la plus courante pour réduire l'indice effectif du guide III-V est de réduire sa largeur en bout de guide. Cette solution est illustrée sur la **figure A-41(a)**, avec une architecture en double tapers inversés [60]. Fabriqué avec des lithographies UV standards, une pointe de 500nm de large a été formée en fin de guide III-V et a permis de coupler la lumière du laser dans un guide silicium de 400nm d'épaisseur, et d'obtenir un laser très performant. Cependant, pour coupler la lumière du laser dans un guide encore plus fin, il serait nécessaire de diminuer encore plus la largeur de la pointe III-V (en-dessous de 300nm pour des guides silicium de 300nm d'épaisseur), ce qui est extrêmement difficile à obtenir de façon répétable avec des moyens standards. De telles pointes ont été réalisées récemment pour un laser hybride III-V sur silicium (voir **figure A-41(b)**), fabriquées avec de la lithographie à faisceau d'électrons et ont permis de coupler la lumière du laser dans un guide de 300nm d'épaisseur [153]. Cependant les performances de ce laser étant très limitées, il est difficile de juger la valeur de cette approche.

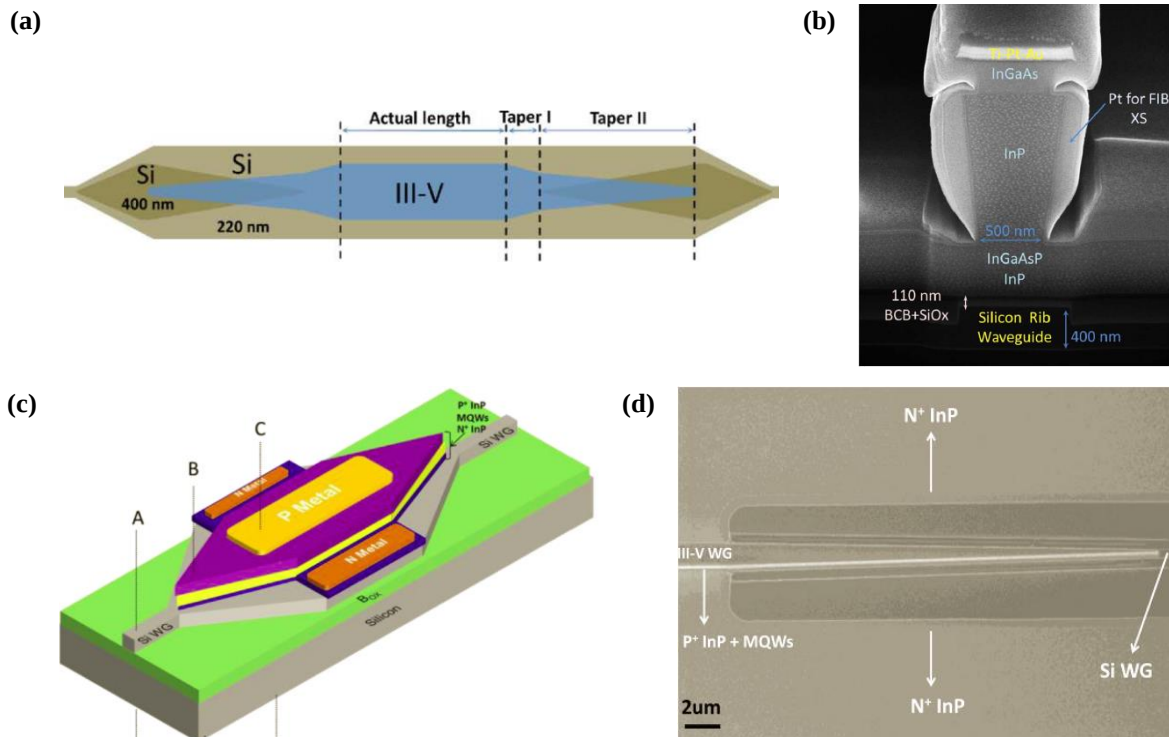


Figure A-41: (a) Vue schématique d'une architecture de double tapers inversés, et (b) image au microscope électronique à balayage de sa pointe de III-V [60]. (c) Vue schématique d'une architecture de double tapers monodirectionnels, et (d) image au microscope électronique à balayage de sa pointe de III-V [153].

Une autre approche consiste à augmenter localement l'épaisseur des guides de silicium utilisés pour le laser, afin d'augmenter leurs indices effectifs. Une solution basée sur une épitaxie de silicium a été proposée afin de former un laser hybride III-V sur silicium [144]. Bien que fonctionnelle, cette solution est limitée en terme d'intégration par le budget thermique élevé de l'épitaxie de silicium. La solution que nous proposons est de déposer une couche de silicium amorphe afin d'atteindre l'épaisseur nécessaire pour le couplage. Contrairement au silicium polycristallin, qui requiert des températures élevées pour réduire son absorption, le silicium amorphe peut être déposé à des températures inférieures à 400°C (le rendant compatible avec nos problématiques d'intégration) tout en offrant des pertes optiques réduites, comparables au silicium monocristallin [162]–[165]. Pour ce faire, les liaisons pendantes (qui sont les sources de pertes supplémentaires du silicium amorphe) doivent être passivées par de l'hydrogène, formant ainsi du a-Si:H. Les guides en a-Si:H voient cependant leurs pertes augmenter exponentiellement avec la température, si chauffés au-delà de 400°C (la limite en température dépendant de l'environnement du recuit [163]).

Afin de valider cette approche, une démonstration est réalisée avec un laser hybride III-V sur silicium dont les dimensions sont similaires à celles utilisées pour le précédent laser, mais avec un guide silicium composé de a-Si:H et de silicium monocristallin (voir **figure A-42(a)**). Bien que le a-Si:H devrait ajouter des pertes supplémentaires (dont la valeur précise n'est pas connue), celles-ci devraient rester limitées, puisque le confinement du supermode pair dans la partie amorphe du guide silicium restera inférieur à 33% (voir **figure A-42(b)**).

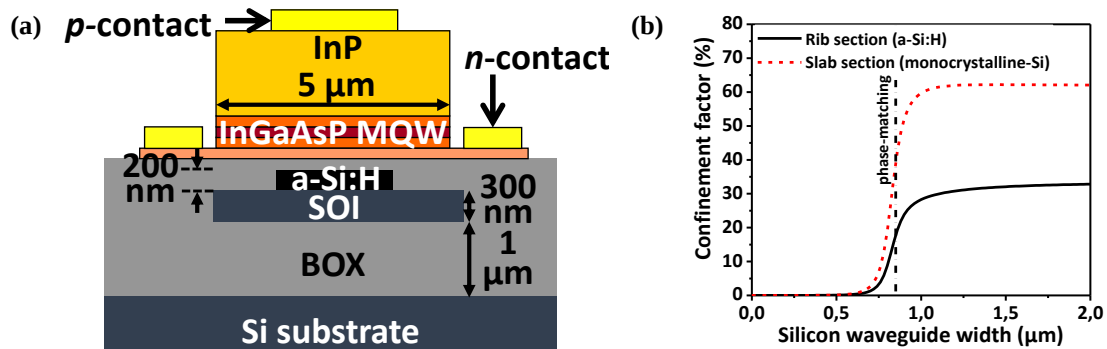


Figure A-42: (a) Vue transverse schématique du laser basé sur un bicouche de a-Si:H et de silicium monocristallin. (b) Facteur de confinement du supermode pair dans l'arête et le slab du guide silicium.

Pour réaliser cette démonstration, une structure de type DFB est choisie pour le laser (puisque'il n'y a pas besoin de contrôler la longueur d'onde d'émission). Contrairement au DBR, le réseau de Bragg utilisé pour former la cavité laser est situé sous le guide III-V (voir **figure A-43**). S'il n'y a pas de défaut dans le réseau, la longueur de la cavité est égale à la moitié d'une période. Dans ce cas, deux modes situés aux bords du spectre de réflexion du réseau vont atteindre le seuil laser simultanément. Afin d'obtenir un laser monomode, il est nécessaire d'intégrer un défaut dans le réseau, qui est généralement introduit en ajoutant une longueur d'une demi-période dans le réseau (c'est le cas de notre structure laser). Cependant, les autres modes peuvent encore remplir les conditions pour atteindre le seuil laser dans certaines conditions.

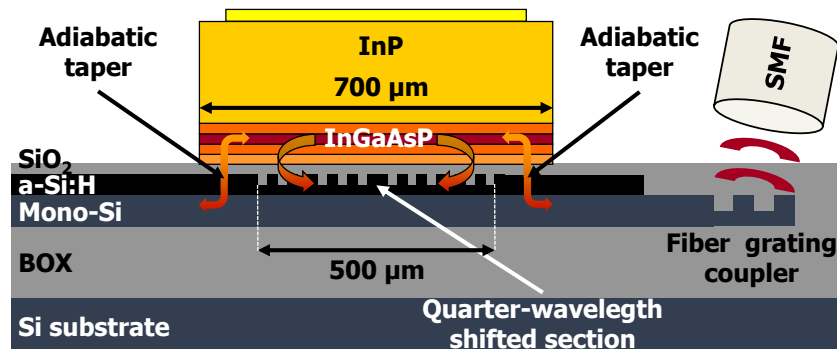


Figure A-43: Vue schématique du laser basé sur un bicouche de a-Si:H et de silicium monocristallin.

L'étape suivante est de fixer les paramètres géométriques du réseau de Bragg. Ici, la longueur du réseau est fixée à $500\mu\text{m}$, et son rapport cyclique à 0.5. Les paramètres restants sont donc la période du réseau, la profondeur de gravure, et la largeur du guide silicium où est défini le réseau. Le réseau étant situé sous le guide III-V, il est nécessaire de limiter la largeur du guide silicium à $0.8\mu\text{m}$, afin de garantir un confinement du supermode pair dans les puits quantiques supérieur à 10% (voir **figure A-30(b)**). Pour évaluer la réflectance du réseau, il est préférable d'évaluer sa constante de couplage $\kappa_g \approx 2\Delta\tilde{n}/\lambda_g$ (où $\Delta\tilde{n}$ est la différence entre les indices effectifs des zones gravées et non-gravées du réseau, et λ_g est la longueur d'onde de Bragg). Celle-ci est liée à la réflectance maximale du réseau par $R_p = \tanh^2(\kappa_g L_g)$, où L_g est la longueur du réseau. Ainsi, pour que la réflectance maximale de chaque côté du laser soit comprise entre 10 et 90%, κ_g doit être compris entre 6 et 72cm^{-1} (voir **figure A-44(a)**). Comme on peut le voir sur la **figure A-43(b)**, cette constante de couplage du réseau dépend principalement de la largeur du guide silicium, et peu de la profondeur de gravure (surtout si celle-ci est au-delà de 60nm). Afin d'obtenir des valeurs suffisamment élevées de constante de couplage du réseau, il est nécessaire d'élargir le guide silicium au-delà de $0.6\mu\text{m}$, afin d'augmenter le confinement du supermode dans le guide silicium. Finalement, la largeur du guide silicium est fixée à $0.8\mu\text{m}$, et la profondeur de gravure à 75nm , afin d'obtenir une constante de couplage du réseau de 40cm^{-1} . Pour cette géométrie de guide, l'indice effectif moyen du supermode pair dans les zones gravées et non-gravées du réseau est proche de 3.3. Pour avoir une longueur d'onde de Bragg de $1.31\mu\text{m}$, la période du réseau est fixée à $\Lambda = \lambda_g/(2\tilde{n}) = 1310/(2 \times 3.3) = 198\text{nm}$. Puisque la largeur du guide silicium sous le guide III-V est fixée par le réseau de Bragg, un taper adiabatique adapté, avec une largeur d'entrée fixée à $0.8\mu\text{m}$, a été conçu avec la méthode décrite dans la section précédente.

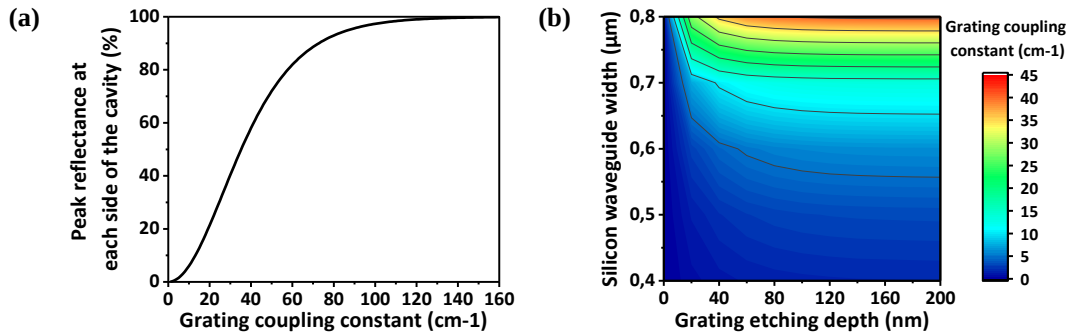


Figure A-44: (a) Réflectance maximale de chaque côté de la cavité en fonction de la constante de couplage du réseau pour une longueur de $500/2 = 250\mu\text{m}$. (b) Constante de couplage du réseau du supermode pair en fonction de la profondeur de gravure du réseau, et de la largeur du guide de silicium sous le guide III-V. La longueur d'onde de Bragg visée en simulation est de $1.31\mu\text{m}$.

Les étapes de fabrication des lasers hybrides III-V sur silicium basé sur un bicouche de a-Si:H et de silicium monocristallin sont détaillées sur la **figure A-45**. Celles-ci sont très proches de celles utilisées pour le transmetteur hybride III-V sur silicium, avec quelques différences notables. La fabrication débute sur des substrats SOI de 200mm de SOITEC, avec une couche de SOI de 300nm et un BOX de $1\mu\text{m}$. Une couche d'oxyde de 20nm est d'abord déposée, suivie d'un dépôt de 220nm de silicium amorphe à 350°C . Bien qu'il eut été préférable de déposer et graver cette couche juste avant le collage des plaques de III-V, cette approche est plus proche de nos procédés de fabrication usuels, et permet de faire subir toutes les étapes de fabrication à la couche de silicium amorphe. Quant à la couche d'oxyde, elle est utilisée comme couche d'arrêt pour la gravure du silicium amorphe. Une fois les différents composants silicium gravés, ceux-ci sont encapsulés et planarisés avant le collage des plaques III-V de 50mm . La composition de ces plaques (fournies par le III-V lab) est proche de celle de la **table A-4**, mais avec un dopage deux fois plus élevé dans les zones dopées p, une couche d'InP dopée p de $1.8\mu\text{m}$, et un pic de photoluminescence à $1.28\mu\text{m}$ pour les puits quantiques. L'épaisseur de collage visée de 100nm est cette fois respectée. Après le collage, la couche sacrificielle d'InP n'est pas retirée juste après le retrait du substrat, mais après une étape d'implantation H^+ , destinée à isoler électriquement les bords du guide III-V. Après le dépôt et la gravure du masque dur SiN utilisé pour graver les guides III-V, les plaques de SOI sont réduites de 200mm à 75mm autour des plaques de III-V, afin de poursuivre leur fabrication dans la salle blanche dédiée aux étapes de fabrication des guides III-V.

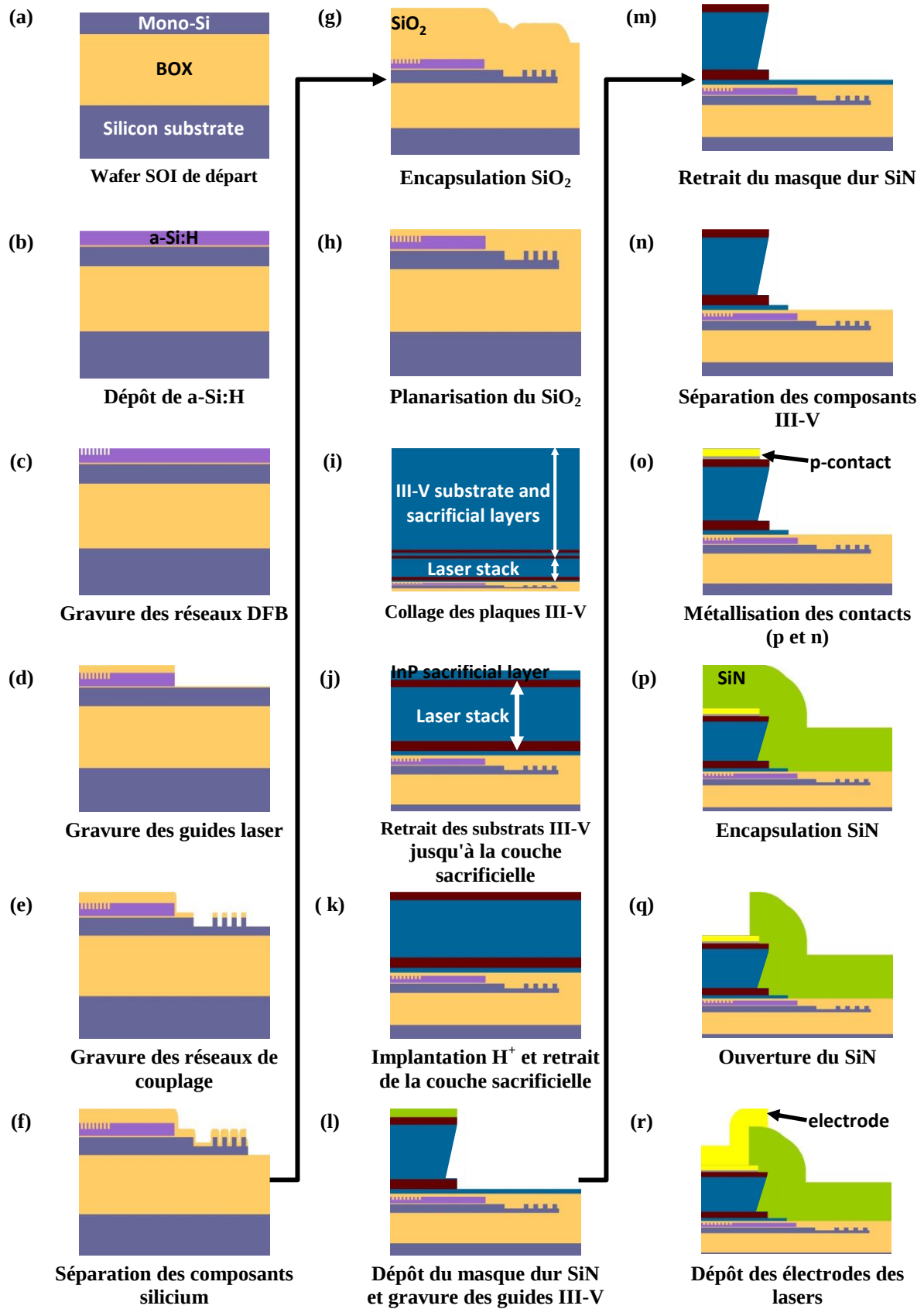


Figure A-45: Etapes de fabrication des transmetteurs hybrides III-V sur silicium basé sur un bicouche de a-Si:H et de silicium monocristallin.

Les guides III-V sont gravés en utilisant la même combinaison de gravures sèches et humides que pour le cas précédent. Après le retrait du masque dur SiN et la séparation des guides III-V, les contacts p et n des lasers sont formés par lift-off. Les structures sont ensuite encapsulées dans une couche de SiN de $2\mu\text{m}$, qui est ouverte au niveau des contacts. Pour finir, les électrodes du laser sont aussi formées par lift-off. Pendant toute la fabrication, un budget thermique de 400°C a été respecté, à l'exception du recuit du contact p du laser (420°C).

Les performances du meilleur des lasers fabriqués sont présentées sur la **figure A-46**. La puissance émise d'un côté du laser est transmise jusqu'à un réseau de couplage (dont les pertes à la longueur d'onde émise par le laser sont évaluées à -5.5dB), et collectée dans une fibre monomode. A 25°C , le courant de seuil du laser est de 50mA , et la puissance maximum couplée dans la fibre est de 4.7mW , pour un courant de 186mA . En prenant en compte la puissance émise des deux côtés du laser, ainsi que les pertes liées au réseau de couplage, la puissance maximale dans le guide silicium est estimée à 33.6mW . Cependant il doit être noté qu'au-delà de 144mA , le laser souffre d'une compétition modale, et perd son caractère monomode. Néanmoins, la « Wall-plug efficiency » qui correspond au ratio de la puissance optique émise sur la puissance électrique injectée dans le laser (évaluée à l'aide de la **figure A-46(b)**), atteint son maximum de 10% pour un courant de 140mA , avant la compétition de modes.

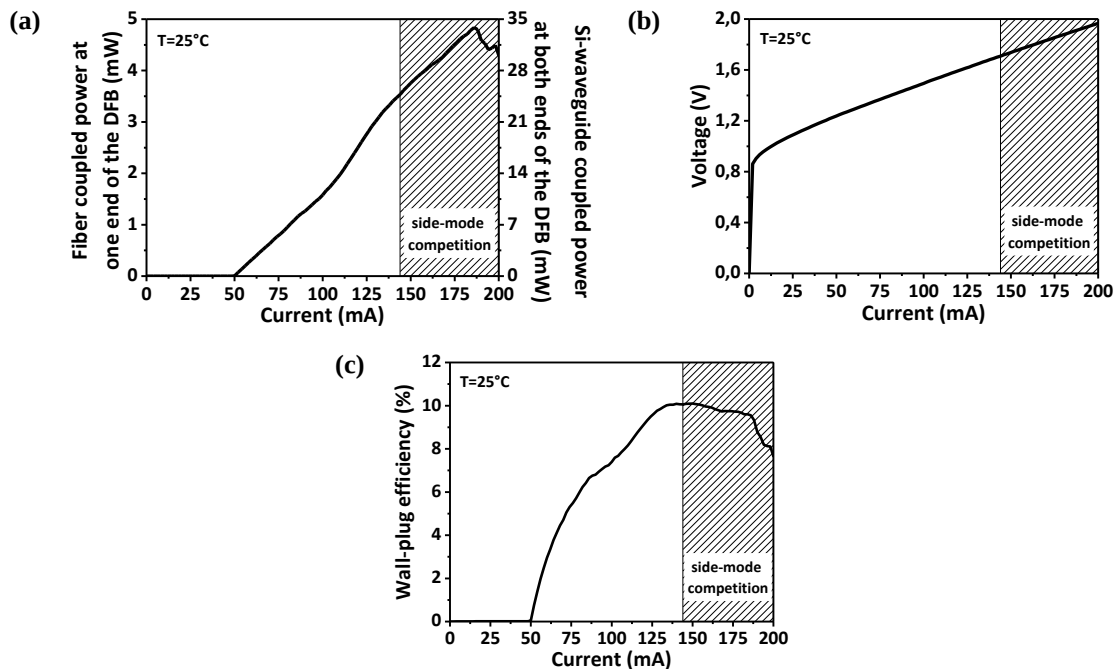


Figure A-46: (a) Puissance optique mesurée à la sortie du réseau de couplage (axe de gauche) et estimée dans le guide silicium (axe de droite), (b) tension entre les électrodes laser, et (c) « Wall-plug efficiency », tous en fonction du courant injecté dans le laser.

Les caractéristiques spectrales du laser sont aussi étudiées et visibles sur la **figure A-47**. Ainsi, les spectres de ce laser (voir **figure A-47(a)** et **(b)**) prouvent qu'il est monomode, avec un SMSR maximum de 51dB . Cependant l'évolution de son SMSR avec le courant (voir **figure A-47(c)**), montre que bien qu'il soit supérieur à 45dB jusqu'à un courant de 144mA , il chute ensuite à $15\text{--}20\text{dB}$. Cette chute de SMSR provient en fait d'une compétition entre la longueur d'onde d'émission du laser (au centre du spectre de réflexion du réseau), et l'un des modes situé au bord du spectre de réflexion du réseau. Cet effet est attribué au « spatial hole burning effect », qui est typique des lasers DFB avec un défaut induit dans leur cavité à un courant d'injection élevé, réduisant leur courant maximum d'utilisation. Néanmoins, la longueur d'onde émise par le laser reste inchangée malgré cette compétition de modes (voir **figure A-47(d)**).

Malgré cette limitation, les performances de ce laser DFB se situent à l'état de l'art des lasers hybrides III-V sur silicium démontrés jusqu'à présent dans la littérature. Cette démonstration prouve ainsi la viabilité de l'utilisation du silicium amorphe pour permettre de coupler efficacement la lumière en sortie d'un laser hybride dans un guide fabriqué à partir d'une fine couche de SOI. De plus, grâce à sa faible température de déposition, cette couche peut être intégrée en fin de fabrication, limitant ainsi son impact sur les composants déjà existants.

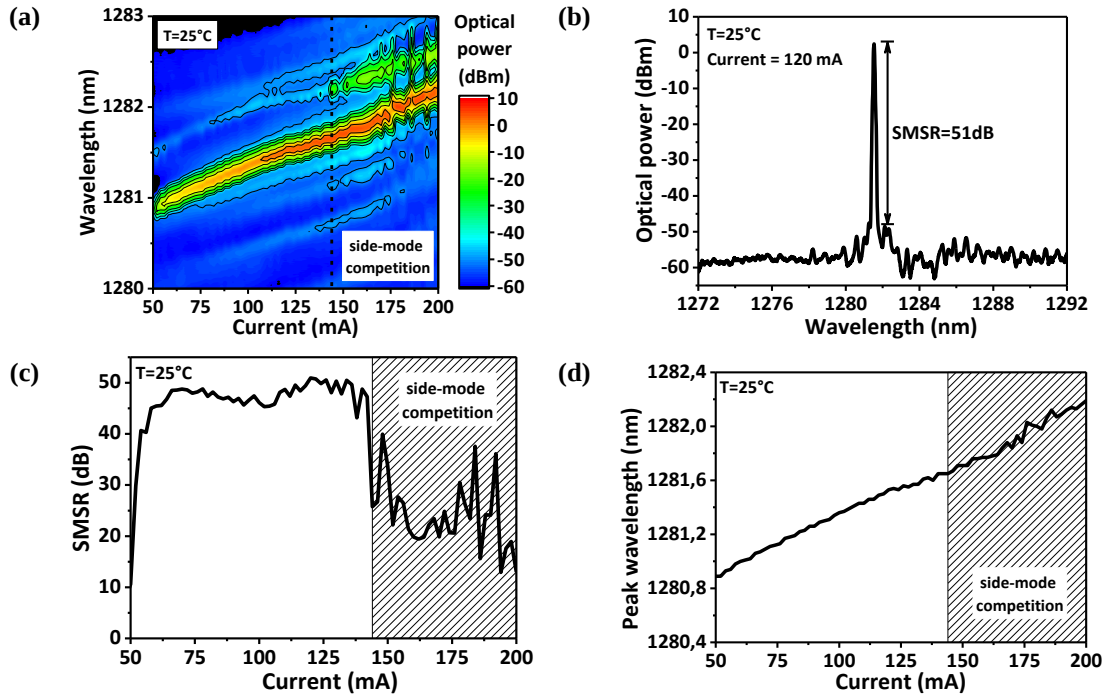


Figure A-47: (a) Spectre du laser en fonction du courant injecté. (b) Spectre du laser à courant fixé (120mA). (c) SMSR et (d) longueur d'onde d'émission en fonction du courant injecté.

A-5. Conclusions et perspectives

L'objectif de ce travail de thèse était de concevoir et produire un transmetteur à haut débit avec une source laser intégrée pour la photonique sur silicium, ciblant des applications telles les communications dans les centres de données. Ce transmetteur, basé sur les spécifications de la norme « 100GBASE-LR4 », devait être aussi compatible que possible avec les plateformes de fabrication de photonique sur silicium actuelles. Ainsi, nous avons réalisé des transmetteurs basés sur la co-intégration de lasers hybrides III-V sur silicium et de modulateurs Mach-Zehnder silicium. La conception des différents composants actifs, ainsi que leur fabrication et leur caractérisation ont été présentées. Les modulateurs silicium ont montré d'excellentes performances à 25Gb/s, bien que limités par les pertes relativement élevées des composants passifs, mais aussi par l'utilisation d'un unique niveau de métallisation, qui a limité les possibilités de conception. Les mesures ont néanmoins montré que l'amélioration des performances des composants passifs permettrait d'atteindre les objectifs visés en termes d'OMA, pour des puissances lasers à la portée des lasers III-V sur silicium actuels. Les transmetteurs hybrides III-V sur silicium fabriqués par la suite ont eux aussi permis d'atteindre des débits de 25Gb/s, pour des distances de 10km., et à deux longueurs d'ondes différentes. Bien que des modulations à d'autres longueurs d'ondes n'ont pas été réalisées, celles-ci auraient pu être effectuées en utilisant les chauffeuses situées au-dessus des réseaux de Bragg des DBR, qui ont montré lors de tests annexes, qu'il était possible d'augmenter la longueur d'onde des lasers de 8.5nm. Bien que fonctionnels, ces transmetteurs sont limités sur plusieurs points à cause de différents incidents durant leur fabrication. La plus importante de ces limitations est la faible puissance de sortie des transmetteurs, causée par la faible puissance émise par les lasers, ainsi que par les pertes élevées des composants passifs formant le transmetteur. Ces pertes, déjà élevées durant les tests des modulateurs seuls, étaient encore accrues après la co-intégration avec les lasers hybrides III-V sur silicium. Une cause possible de cette augmentation pourrait être attribuée à des procédés de gravure additionnels, les composants autres que ceux dédiés au laser et basés sur des guides silicium de 300nm d'épaisseur ayant été formés à partir d'une couche de SOI de 500nm. En effet, cette épaisseur de 500nm est utilisée pour former les composants nécessaires au couplage de la lumière générée par les guides de III-V dans les guides silicium. Cette différence d'épaisseur entre les composants dédiés au laser (qui requièrent une épaisseur de silicium de 500nm), et ceux formant le reste du circuit silicium (qui sont basés sur des épaisseurs de SOI inférieures à 300nm dans la majorité des

plateformes standards de fabrication de photonique sur silicium) est par ailleurs un frein à l'intégration des lasers hybrides III-V sur silicium. Par conséquent, une solution basée sur la déposition d'une couche silicium amorphe pour augmenter localement l'épaisseur des guides silicium a été proposée, afin de coupler la lumière générée par les guides III-V dans des guides silicium formés à partir de SOI de 300nm d'épaisseur. De plus, pouvant être déposée à des températures inférieures à 400°C , cette couche de silicium amorphe peut être intégrée en fin de fabrication, et ainsi permettre de ne pas influencer les composants standards déjà présents. Une réalisation de lasers DFB basés sur cette approche a permis d'obtenir d'excellents résultats proches de l'état de l'art, et de démontrer la viabilité de cette approche. Ainsi, malgré les différentes limitations des différents composants fabriqués, ce travail de thèse prouve que la technologie photonique sur silicium actuelle, associée avec l'intégration hétérogène de matériaux III-V sur silicium permet de fabriquer des systèmes intégrés compatibles avec les contraintes de la transmission de données à plusieurs kilomètres de distance.

Concernant les nombreuses suites possibles à donner à ce travail, la plus importante à court terme serait de fabriquer entièrement ces transmetteurs sur des plaques de SOI de 200mm ou même de 300mm , contrairement aux transmetteurs réalisés dans cette thèse, qui ont nécessité de réduire le diamètre des plaques à 75mm pour terminer leur fabrication dans une salle blanche différente. Pour ce faire, plusieurs problèmes doivent encore être résolus. Pour commencer, il est nécessaire d'améliorer le rendement de collage des puces de III-V, au lieu de plaques complètes de III-V, puisqu'une grande partie du III-V est non-utilisée pendant la fabrication. Ensuite, pour éviter la réduction des plaques de SOI et le transfert dans une autre salle blanche, il faudrait résoudre les problèmes de contamination liés à l'utilisation des matériaux III-V et des métallisations du laser à base d'or. Pour finir, bien qu'il serait alors possible de fabriquer des lasers hybrides III-V sur silicium seuls sur des plaques de 200mm ou 300mm , les guides III-V utilisés restent trop épais pour être intégrés avec les différents niveaux de métallisations que l'on peut trouver dans les plateformes standards de fabrication de photonique sur silicium, dont l'épaisseur totale est du même ordre de grandeur. La réalisation de ces niveaux de métal nécessitant plusieurs étapes de planarisation, ils ne peuvent pas être intégrés après les guides III-V, qui eux-mêmes nécessitent d'être proches des guides silicium. Une solution proposée par un autre étudiant en doctorat (Jocelyn Durel) est d'intégrer le guide III-V sur la face arrière du circuit photonique (aussi appelée intégration laser « back-side »). Dans cette approche, une fois les composants silicium fabriqués avec leurs niveaux de métallisation, une plaque de silicium est collée au-dessus de la plaque SOI, avant de retourner l'ensemble. Le substrat, ainsi que le BOX du wafer SOI sont alors retirés, avant de coller les matériaux III-V et de former les guides lasers (voir **figure A-48**). Dans ce cas de figure, la lumière est aussi collectée par la face arrière du circuit photonique, à l'aide d'un miroir métallique situé sous le réseau de couplage. Cette approche permet non seulement de conserver les multiples niveaux de métal d'un circuit photonique standard (ce qui bénéficierait à la conception des modulateurs silicium), mais aussi de repousser les étapes de fabrication des sources lasers à la fin de la fabrication du circuit, ce qui permettrait de simplifier l'intégration des matériaux III-V et des métallisations à base d'or dans les zones « far back-end » des salles blanches. Pour finir, la lumière étant collectée par la face arrière du circuit photonique, il serait alors possible de coller un circuit électronique de contrôle sur la face avant.

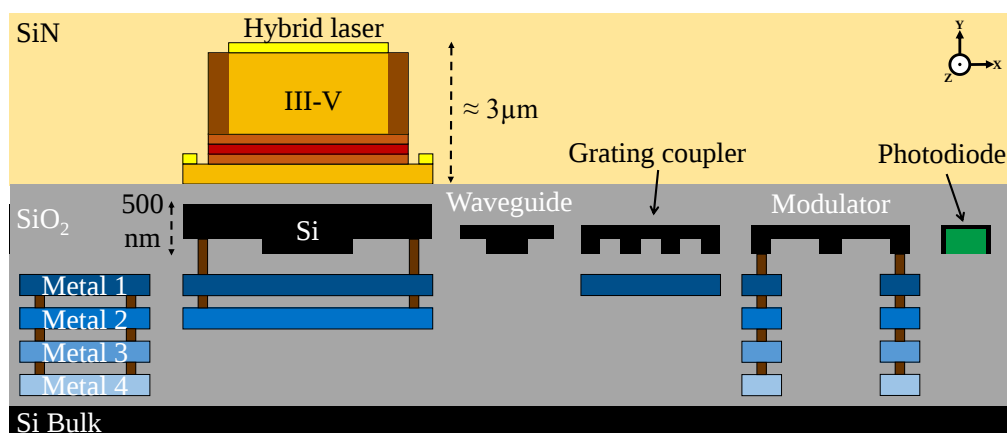


Figure A-48: Vue schématique de l'intégration laser "back-side" [170].

AUTORISATION DE SOUTENANCE

Vu les dispositions de l'arrêté du 7 août 2006,

Vu la demande du Directeur de Thèse

Monsieur C. SEASSAL

et les rapports de

Mme D. MARRIS-MORINI

Professeur - Centre de Nanosciences et de Nanotechnologies (C2N) - site d'Orsay - UMR9001
Bât. 220 - rue André Ampère - Université Paris-Sud - 91405 ORSAY cedex

et de

M. G. MORTHIER

Professeur - Ghent University - IMEC - Department of Information Technology - Office 140.018
Technologiepark-Zwijnaarde 15 - 9052 Gent - Belgique

Monsieur FERROTTI Thomas

est autorisé à soutenir une thèse pour l'obtention du grade de **DOCTEUR**

Ecole doctorale ELECTRONIQUE, ELECTROTECHNIQUE, AUTOMATIQUE

Fait à Ecully, le 30 novembre 2016

P/Le directeur de l'E.C.L.
La directrice des Etudes



Design, fabrication and characterization of a hybrid III-V on silicon transmitter for high-speed communications.

Abstract: For several years, the volume of digital data exchanged across the world has increased relentlessly. To manage this large amount of information, high data transmission rates over long distances are essential. Since copper-based interconnections cannot follow this tendency, high-speed optical transmission systems are required in the data centers. In this context, silicon photonics is seen as a way to obtain fully integrated photonic circuits at an expected low cost. While this technology has experienced significant growth in the last decade, the high-speed transmitters demonstrated up to now are mostly based on external laser sources. Thus, the aim of this PhD thesis was to design and produce a high-speed silicon photonic transmitter with an integrated laser source.

This transmitter is composed of a high-speed silicon Mach-Zehnder, co-integrated on the same wafer with a hybrid III-V on silicon distributed Bragg reflector laser, whose emission wavelength can be electrically tuned in the 1.3 μ m wavelength region. The design of the various elements constituting both the laser (III-V to silicon adiabatic couplers, Bragg reflectors) and the modulator (p - n junctions, travelling-wave electrodes) is thoroughly detailed, as well as their fabrication. During the characterization of the transmitters, high-speed data transmission rates up to 25Gb/s, for distances up to 10km are successfully demonstrated, with the possibility to tune the operating wavelength up to 8.5nm. Additionally, in order to further improve the integration of the laser source with the silicon photonic circuit, a solution based on the low-temperature (below 400°C) deposition of an amorphous silicon layer during the fabrication process is also evaluated. Tests on a distributed feed-back laser structure have shown performances at the state-of-the-art level (with output powers above 30mW), thus establishing the viability of this approach.

Keywords: Silicon photonics, integrated photonic circuits, hybrid III-V on silicon lasers, Mach-Zehnder modulators.

Conception, fabrication et caractérisation d'un transmetteur hybride III-V sur silicium pour communications à haut débit.

Résumé: Depuis plusieurs années, le volume de données échangé à travers le monde augmente sans cesse. Pour gérer cette large quantité d'information, des débits élevés de transmission de données sur de longues distances sont essentiels. Puisque les interconnexions à base de cuivre ne peuvent pas suivre cette tendance, des systèmes de transmission optique rapides sont requis dans les centres de données. Dans ce contexte, la photonique sur silicium est considérée comme une solution pour obtenir des circuits photoniques intégrés à un coût réduit. Bien que cette technologie ait connu une croissance significative au cours de la dernière décennie, les transmetteurs actuels à haut débit de transmission sont principalement basés sur des sources laser externes. Par conséquent, l'objectif de ce travail de thèse était de concevoir et produire un transmetteur à haut débit de transmission de données pour la photonique sur silicium, doté d'une source laser intégrée.

Ce transmetteur se compose d'un modulateur silicium de type Mach-Zehnder, co-intégré sur la même plaque avec un laser hybride III-V sur silicium à réseaux de Bragg distribués, dont la longueur d'onde d'émission peut être contrôlée électriquement autour de 1.3 μ m. La conception des différents éléments constituant à la fois le laser (coupleurs adiabatique entre le III-V et le silicium, miroirs de Bragg) et le modulateur (jonctions p - n , électrodes à ondes progressives) est détaillée, de même que leur fabrication. Pendant la caractérisation des transmetteurs, des taux de transmission de données jusqu'à 25Gb/s, pour des distances allant jusqu'à 10km ont été démontrés avec succès, avec la possibilité de contrôler la longueur d'onde jusqu'à 8.5nm. Par ailleurs, afin d'améliorer l'intégration de la source laser avec le circuit photonique sur silicium, une solution basée sur le dépôt à basse température (en-dessous de 400°C) d'une couche de silicium amorphe pendant la fabrication est aussi évaluée. Des tests sur une cavité laser à contre-réaction distribuée ont montré des performances au niveau de l'état de l'art (avec des puissances de sortie supérieures à 30mW), prouvant ainsi la viabilité de cette approche.

Mots-clés: Photonique sur silicium, circuits photoniques intégrés, lasers hybrides III-V sur silicium, modulateurs Mach-Zehnder.