

Silicon Photonic Microring Resonators: A Comprehensive Design-Space Exploration and Optimization under Fabrication-Process Variations

Asif Mirza, *Student Member, IEEE*, Febin Sunny, *Student Member, IEEE*, Peter Walsh,
Karim Hassan, Sudeep Pasricha, *Senior Member, IEEE*, and Mahdi Nikdast, *Senior Member, IEEE*

Abstract—Silicon photonic microring resonators (MRRs) offer many advantages (e.g., compactness) and are often considered as the fundamental building block in optical interconnects and emerging photonic nanoprocessors and accelerators. Such devices are, however, sensitive to inevitable fabrication-process variations (FPVs) stemming from optical lithography imperfections. Consequently, silicon photonic integrated circuits (PICs) integrating MRRs often suffer from high power overhead required to compensate for the impact of FPVs on MRRs, and hence realizing a reliable operation. On the other hand, the design space of MRRs is complex, including several correlated design parameters, thereby further exacerbating design optimization of MRRs under FPVs. In this paper, we present, for the first time, a comprehensive design-space exploration in passive and active MRRs under FPVs. In addition, we present design optimization in MRRs under FPVs while considering different performance metrics such as tolerance to FPVs, Quality factor, and 3dB bandwidth in MRRs. Simulation and fabrication results obtained by measuring multiple fabricated MRRs designed using our design-space exploration demonstrate a significant 70% improvement on average in the MRRs' tolerance to different FPVs. Furthermore, we apply the proposed design optimization to a case study of a wavelength-selective MRR-based demultiplexer where we show considerable channel-spacing accuracy within 0.5 nm even when the MRRs are placed 500 μm apart on a chip. Such improvements indicate the efficiency of the proposed design-space exploration and optimization to enable power-efficient and variation-resilient PICs and optical interconnects integrating MRRs.

Index Terms—Silicon photonics, microring resonators, design-space exploration, fabrication-process variations.

I. INTRODUCTION

SILICON-ON-INSULATOR (SOI) platform has enabled the implementation of sub-micron optical waveguides and dense packing of various optical functions on a single chip. Among different SOI-based silicon photonic devices designed for photonic integrated circuits (PICs), microring resonators (MRRs) have attracted much attention because of their compact footprint (e.g., radius $\approx 5 \mu\text{m}$) and ability to perform a variety of functions such as optical filtering [1], modulation [2], and spatial switching [3], [4]. The central optical frequency (i.e., resonant wavelength) in MRRs is an

essential parameter that is determined by several key factors in the MRR design space, including waveguide width, SOI thickness, and MRR radius [5]. For dense wavelength-division multiplexing (DWDM) applications, where a large number of optical channels (i.e., wavelengths) are narrowly spaced to boost the bandwidth performance in PICs [6], it is critical to maintain accurate channel spacing and match the central resonant wavelengths of different MRRs to achieve a reliable communication [7]. Nevertheless, the resonant wavelength in an MRR is highly sensitive to the variations in the critical dimensions of the MRR due to inevitable fabrication-process variations (FPVs).

FPVs originate in optical-lithography process imperfections, contributing to different variations in the waveguide thickness and linewidth, waveguide edge roughness and sidewall slope, dopant, etc. [8]. In PICs employing MRRs, FPVs impose resonant-wavelength deviations, leading to severe PIC performance degradation, or in the worst-case, a total circuit failure [9]. For example, prior work shows ≈ 2 nm shift in the resonant wavelength of an MRR with only a single nanometer change in its waveguide thickness [10]. Such a small resonant-wavelength shift is of critical concern in DWDM systems with a large number of MRRs, each tuned to a specific optical channel (i.e., wavelength), with a channel spacing as small as 0.8 nm [11]. Several methods have been proposed to compensate for the impact of FPVs in PICs at run-time by leveraging the thermo-optic [12] and electro-optic effects [13] of silicon. Nevertheless, such methods lead to considerable increase in power dissipation in PICs [14]. For example, internal integrated heaters can be used in MRRs to tune the resonant wavelength by 40 mW/FSR [15], where FSR is the free-spectral range. This rather small power consumption quickly adds up in PICs that integrate a large number of MRRs, where each MRR's resonant wavelength may need to be adjusted by several nanometers (e.g., 9 nm in [16]). This calls for efficient design solutions to minimize the effect of FPVs at design-time in the current fabless silicon photonic ecosystem [10], thus improving PIC tolerance to FPVs and reducing required tuning power consumption after the fabrication. Note that post-fabrication trimming of MRRs is also possible [17], but it will significantly increase the fabrication cost.

The novel contribution of this paper is on developing the first comprehensive design-space exploration for silicon photonic MRRs under different FPVs. In particular, 1) we analyze and model the impact of different FPVs in waveguide

This work was supported by the National Science Foundation (NSF) under grant numbers CCF-1813370, CCF-2006788, and CNS-2046226.

A. Mirza, F. Sunny, P. Walsh, S. Pasricha, and M. Nikdast are with the Department of Electrical and Computer Engineering, Colorado State University, Fort Collins, CO 80523, USA (email: {asifmirz, febins, p8walsh, sudeep, mnikdast}@colostate.edu).

K. Hassan is with University Grenoble Alpes, CEA, Leti, F-38000 Grenoble, France (email: karim.hassan@cea.fr).

width, SOI thickness, slab thickness, and MRR radius (see Fig. 1) on the resonant wavelength of passive and active MRRs, identifying the impact of each variation on the resonant wavelength in such MRRs; 2) we present computationally efficient analytical models that capture the impact of such variations at the physical layer on the MRRs' device-layer performance, including coupling efficiency, Quality factor (Q-factor), and 3dB bandwidth; 3) leveraging our detailed design-space exploration, we develop MRR design optimization to find optimal design parameters in MRRs under different FPVs with a goal of enhancing not only the device tolerance to FPVs but also improving other important factors such as Q-factor and 3dB bandwidth; and, 4) we develop realistic virtual FPV wafer maps accounting for both correlated and radial variations in silicon photonics to further enhance the MRR design-space exploration and optimization.

Simulation results obtained using accurate finite-difference time-domain (FDTD) and eigenmode (FDE) simulations [18] indicate the efficiency of our proposed design-space exploration and optimization. Moreover, our study includes the design and fabrication of multiple MRRs designed using the proposed design-space exploration. Experimental results from measuring several fabricated MRRs demonstrate a good agreement with our simulation results while showing an average 70% improvement in MRRs tolerance to different FPVs. As a case study and by employing correlated virtual FPV wafer maps and actual device layout information, we apply our design optimization to enhance the inter-device matching (i.e., channel-spacing accuracy) in a wavelength-selective MRR-based demultiplexer under different FPVs. We show significant channel-spacing accuracy within 0.5 nm even when the MRRs are positioned far apart at 500 μm on a chip. By minimizing channel-spacing variations, we can compensate for resonance shifts by collectively tuning all the MRRs, hence simplifying the tuning and improving its efficiency. Note that thermal variations also deviate the resonant wavelength in MRRs, but we focus on the impact of FPVs in this paper.

The rest of the paper is organized as follows. Section II summarizes related work on FPV compensation and analysis in PICs. We present efficient analytical models to study MRRs under FPVs in Section III. Section IV includes our detailed design-space exploration in MRRs under different FPVs and also our fabrication results. We present our MRR design optimization solution in Section V, and the case study of a wavelength-selective MRR-based demultiplexer in Section VI. Finally, Section VII concludes our work.

II. RELATED WORK

FPVs in silicon photonics have been studied at different scales, including exploring within-die, die-to-die, wafer-to-wafer, and lot-to-lot variations [19]. In [20], R. G. Beausoleil *et al.* studied the impact of FPVs in identically designed MRRs and reported within-die and within-wafer variations with a variance of 0.5 and 2 nm, respectively. A. V. Krishnamoorthy *et al.* explored the manufacturing tolerance of over 500 four-channel MRRs fabricated in a commercial 130 nm CMOS foundry with 193 nm lithography [21]. Results from characterizing multiple reticles, wafers, and fabrication lots show

that the absolute resonant wavelengths of individual MRRs cannot be controlled across wafers or even across reticles or fields within a wafer. In [22], X. Chen *et al.* studied FPVs in MRRs and racetrack resonators, all identically designed but fabricated in two different establishments (Leti and IMEC). They reported resonant-wavelength variations with a variance of 1.3 nm²/cm in both the MRRs and racetrack resonators.

Through various examples of wavelength filters, W. Bogaerts *et al.* in [10] showed that variability modeling can guide the design process and make circuits more robust against different FPVs. Their work analyzed the performance of a Mach-Zehnder lattice filter under waveguide width and SOI thickness variations to improve the fabrication yield by optimizing circuit layout. In [23], Z. Lu *et al.* proposed a method to characterize FPVs and predict circuit performance under the impact of correlated FPVs. Using this method, they measured the spectral responses of 2074 identical racetrack resonators on a 200 mm wafer fabricated using a 248 nm deep UV (DUV) lithography photonics process. Moreover, for such a fabrication process, they reported standard deviations for waveguide width and SOI thickness of 3.9 and 1.3 nm, respectively. Y. Wang *et al.* in [24] proposed a hierarchical approach that takes into account domain specific knowledge, spatial frequency analysis, and low-rank tensor factorization methods to decompose FPVs into wafer-level, intra-die, and inter-die components.

Similar to this paper, the work in [25]–[27] aims at modeling and improving MRRs under FPVs. Y. London *et al.* in [25] proposed a behavioral model for directional couplers with non-identical waveguide widths, and not for the full design space of MRRs, which is a more complex problem to address. Y. Luo *et al.* in [26] developed an MRR modulator design based on using multi-mode waveguides that can improve accuracy and repeatability of the MRR resonant wavelength. In their work, the MRR is based on an unequal design of ring and bus waveguide widths where the ring waveguide width is increased to reduce phase errors associated with side-wall etch. Z. Su *et al.* in [27] proposed the use of adiabatic rings to design MRRs with high Q-factor and improved tolerance to FPVs, but the designed MRRs are too small (radii of 2–3 μm) and suffer from high optical power losses. Unlike the contributions in this paper, [25]–[27] lacks comprehensive modeling and optimization of MRRs under FPVs, thereby focusing on the experimental design and improvement of the individual MRR designs considered in these papers. Consequently, the results in [25]–[27] are limited to the MRR structures considered in these works. As we will show, our proposed design-space exploration and optimization in this paper can be applied to any MRR design problem, including those in [25]–[27], providing designers with a comprehensive framework to explore and optimize any MRR under FPVs.

Our prior work in [28], [29] studied the impact of FPVs in waveguide width and SOI thickness in passive MRRs. However, we did not study the design space of active MRRs and variations in the MRR radius. In this work, we not only provide complete details of our design-space exploration and optimization for MRRs under different FPVs, we significantly extend our prior work by developing a comprehensive study

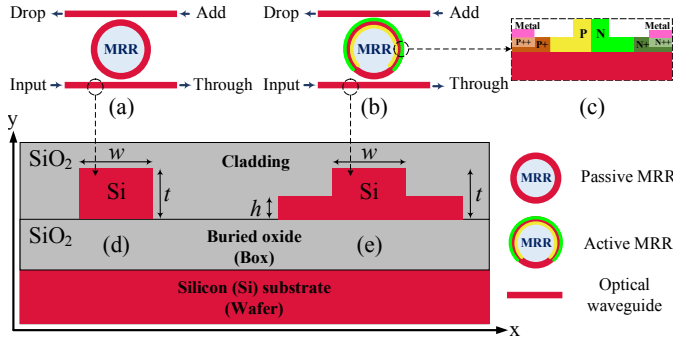


Fig. 1: An overview of an MRR add-drop filter showing waveguide width (w), SOI thickness (t), and slab thickness (h) in (a) a passive and (b) an active MRR with (c) a P-N junction. Cross section of (d) a strip and (e) a ridge waveguide. Here, Si and SiO₂ denote silicon and silicon dioxide, respectively.

and optimization of both passive and active MRRs under FPVs in waveguide width, SOI thickness, slab thickness, and MRR radius. Different from [28], [29], our design-space exploration and optimization in this work includes the complete design space of MRRs (i.e., input/drop waveguide and ring) while considering bending and propagation losses for more accurate analyses. In addition to tolerance to FPVs, our design-space exploration in this work includes several critical factors in the design of passive and active MRRs, including coupling efficiency, Q-factor, and 3dB bandwidth, based on which a set of optimal design parameters are identified for the resilient and high-performance design of MRRs.

III. MODELING SILICON PHOTONIC MRRS UNDER FABRICATION-PROCESS VARIATIONS

In this section, we present the fundamental analytical models required to study the impact of FPVs in MRRs. These models lay the foundation for our proposed design-space exploration and optimization in Sections IV and V. As shown in Figs. 1(a) and 1(b), we consider passive MRRs designed using strip waveguides and active MRRs constructed using ridge waveguides, which allow for electrical connections to be made to the waveguides [7] (e.g., through P-N junctions), as shown in Fig. 1(c). Moreover, Figs. 1(d) and 1(e) show a cross section of a strip and a ridge waveguide while specifying the waveguide width (w), SOI thickness (t), and slab thickness (h). Hereafter and unless specifically mentioned, we use "MRRs" to refer to both passive and active MRRs.

The effective index (n_{eff}) in a waveguide depends on the optical wavelength and the critical dimensions of the waveguide (i.e., w , t , and h —in the case of a ridge waveguide) [5]. The relation can be written as:

$$n_{eff}(\lambda, w, t, h) = \left(\frac{\lambda}{2\pi} \right) \beta(\lambda, w, t, h), \quad (1)$$

where β is the propagation constant and λ is the optical wavelength. Note that numerical solutions of β can be obtained using various numerical and approximation methods (e.g., effective-index method [30]). Leveraging (1), we can define the group index (n_g) in a waveguide as:

$$n_g(\lambda, w, t, h) = n_{eff}(\lambda, w, t, h) - \lambda \frac{\partial n_{eff}(\lambda, w, t, h)}{\partial \lambda}. \quad (2)$$

An MRR can be either on or off resonance depending on the resonant wavelength of the MRR and the optical wavelength on the input waveguide. In an MRR add-drop filter (see Fig. 1(a) as an example), when the round-trip optical phase in the MRR is an integer multiple of 2π , the MRR is on resonance and it drops the input signal. Otherwise, the MRR is off resonance and the optical signal on the input waveguide passes the MRR towards the through port. The resonant wavelength (λ_R) in an MRR can be calculated based on the effective index (defined in (1)) as:

$$\lambda_R(w, t, h, R) = \left(\frac{2\pi R}{m} \right) n_{eff}(\lambda_R, w, t, h), \quad (3)$$

where R is the radius of the MRR and m is an integer that denotes the order of the resonant mode. According to (3), the resonant wavelength in an MRR depends on the critical dimensions of the MRR, including the waveguide width, SOI thickness, slab thickness (in case of an active MRR), and MRR radius, in which any slight variations will deviate the resonant wavelength. Such a deviation is known as the resonant-wavelength shift. Considering the first-order approximation of the waveguide dispersion, the resonant-wavelength shift ($\Delta\lambda_R$) in an MRR can be modeled as:

$$\Delta\lambda_R(\lambda_{R0}, w', t', h', R') = \frac{\lambda_{R0} |n_{eff}(\lambda_{R0}, w, t, h) - n_{eff}(\lambda_{R0}, w', t', h')|}{n_g(\lambda_{R0}, w, t, h)} + \left(\frac{\lambda_{R0} R'}{R} \right) \frac{n_{eff}(\lambda_{R0}, w, t, h)}{n_g(\lambda_{R0}, w, t, h)}. \quad (4)$$

Here, λ_{R0} is the nominal resonant wavelength. Moreover, $w' = w \pm \nu_w$, $t' = t \pm \nu_t$, $h' = h \pm \nu_h$, and $R' = R \pm \nu_R$, where ν_w , ν_t , ν_h , and ν_R are the variations in the waveguide width, SOI thickness, slab thickness, and MRR radius, respectively. In this paper, we assume independent variations (e.g., ν_R is independent of ν_w). Leveraging (4), Fig. 2(a) shows the resonant-wavelength shift in an active MRR while considering one FPV at a time and a variation range of $[-10, 10]$ nm (the x-axis), considered as an example. As can be seen, the resonant wavelength changes almost linearly under all the different variations, and the impact of each variation on the MRR resonant wavelength is different. A similar trend was reported in [21]. As an example, Fig. 2(b) shows the optical spectrum of the MRR in Fig. 2(a)—simulated using the transfer-matrix method [31]—where there is a red and a blue shift with $\nu_w = 5$ nm and $\nu_w = -5$ nm, respectively. There is a good agreement between the resonant-wavelength shift results calculated in Fig. 2(a) and simulated in Fig. 2(b).

FPVs also impact the through- and drop-port response in MRRs, introducing extra power losses when an optical signal passes or drops into an MRR [14]. Such power loss will degrade the power efficiency in PICs employing MRRs. Here, we analyze the impact of FPVs on the coupling strength in the coupling region in an MRR (i.e., between the input/drop waveguide and the ring waveguide—see Fig. 3) that determines the through- and drop-port response in the MRR. We define κ as the cross-over coupling coefficient between the input/drop waveguide and the ring, and s as the straight-through coefficient associated with the power that remains

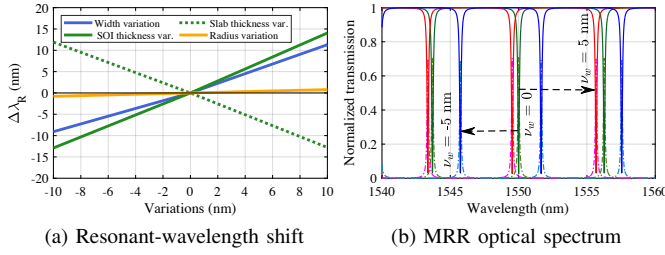


Fig. 2: (a) Resonant-wavelength shift ($\Delta\lambda_R$) in an active MRR calculated using (4) with $w = 400$ nm, $t = 220$ nm, $h = 90$ nm, and $R = 10$ μ m under variations in the waveguide width, SOI thickness, slab thickness, and radius (x -axis). (b) Optical spectrum of the MRR simulated using transfer-matrix method [31] with $\lambda_{R0} = 1550$ nm where the resonant wavelength shifts because of width variations ($\nu_w = \pm 5$ nm).

on the input waveguide [5]. We assume that the input and drop waveguides are coupled symmetrically to the ring, and consider a lossless coupler where $|\kappa|^2 + |s|^2 = 1$. Both κ and s can be calculated precisely based on accurate numerical methods (e.g., FDTD) or using approximate methods such as the supermode theory [14]. A compact systematic model to study the impact of FPVs on the cross-over coupling coefficient (κ) in MRRs can be defined as:

$$\kappa(\lambda_R, w', t', h', R') = f(n_{e/o}(\lambda_R, w', t', h'), R', g^{-1}). \quad (5)$$

Considering (5), κ is a function of the effective indices of the even (n_e) and odd (n_o) supermodes in the coupling region of an MRR that change under FPVs [32] (see Appendix A). Moreover, κ is directly proportional to the MRR radius (R) and inversely proportional to the gap (g), hence g^{-1} in (5). The gap is defined as the edge-to-edge distance between the input/drop waveguide and the ring waveguide. Note that $f()$ in (5) is a function to calculate the cross-over coupling in MRRs that can be defined based on the method described in [7].

Q-factor is the measure of the sharpness of the resonance relative to its central frequency that impacts the optical channel spacing, crosstalk, bandwidth, and other characteristics in MRRs [5]. In particular, it is desirable to have MRRs with a high Q-factor for DWDM applications. Nevertheless, FPVs can deteriorate the Q-factor in MRRs. Assuming a lossless coupler and using (2) and (5), the Q-factor in an MRR add-drop filter under FPVs can be defined as:

$$Q(\lambda_R, w', t', h', R') = \left(\frac{\pi L}{\lambda_R} \right) \frac{n_g(\lambda_R, w', t', h') \sqrt{a(1 - \kappa^2(\lambda_R, w', t', h', R'))}}{1 - a(1 - \kappa^2(\lambda_R, w', t', h', R'))}. \quad (6)$$

Here, $L = 2\pi R$ is the round-trip length in the MRR. Also, a is the single-pass amplitude transmission, which includes the propagation loss in the MRR and loss in the couplers, and can be calculated as $a^2 = \exp(-\alpha L)$, where α is the power attenuation coefficient [5].

The optical bandwidth of interest plays an important role in deciding the optimal design parameters in MRRs. The optical bandwidth in an MRR can be defined based on the frequency at which half the power is incident in the channel (i.e., frequency

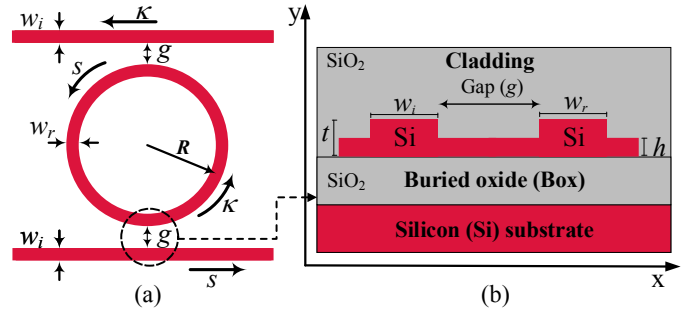


Fig. 3: A cross section of the coupling region in an active MRR add-drop filter with different physical- and device-level design parameters. Here, w_i is the input/drop waveguide width and w_r denotes the ring waveguide width. Note that $h = 0$ for a passive MRR.

at 3dB). Therefore, employing (2) and (5), the 3dB bandwidth in an MRR can be attained by observing the full width at half maximum (FWHM) of the resonance spectrum, defined as:

$$FWHM(\lambda_R, w', t', h', R') = \left(\frac{\lambda_R^2}{\pi L} \right) \frac{1 - a(1 - \kappa^2(\lambda_R, w', t', h', R'))}{n_g(\lambda_R, w', t', h') \sqrt{a(1 - \kappa^2(\lambda_R, w', t', h', R'))}}. \quad (7)$$

The 3dB bandwidth of an MRR used in a DWDM-based add-drop configuration should be large enough to accommodate the bandwidth of the optical signal to be dropped. A narrow bandwidth will result in undesired truncation of the signal spectrum causing distortions [33]. Nonetheless, a large 3dB bandwidth should be avoided too as it can cause severe crosstalk noise (e.g., inter-channel crosstalk) if the channel density is high [34]. Leveraging (7) and considering c to be the speed of light in vacuum, the 3dB bandwidth in an MRR ($\Delta\nu$, in GHz) can be modeled as [7]:

$$\Delta\nu(\lambda_R, w', t', h', R') = \left(\frac{c}{\lambda_R^2} \right) FWHM(\lambda_R, w', t', h', R'). \quad (8)$$

The analytical models proposed in this section are computationally inexpensive and derived to capture the impact of various physical-level FPVs on the critical dimensions and device-level performance in MRRs. Such models can enable a design-space exploration in MRRs under different variations, as discussed in the next section.

IV. MRR DESIGN-SPACE EXPLORATION UNDER FABRICATION-PROCESS VARIATIONS

Leveraging the proposed analytical models in Section III, we present a comprehensive design-space exploration for MRRs under different FPVs in this section. In particular, we analyze the impact of different variations on the resonant wavelength, cross-over coupling, Q-factor, and 3dB bandwidth in MRRs. In addition, we explore the impact of changing the design parameters in MRRs (e.g., the waveguide width) on the device performance under different FPVs. As the input/drop and ring waveguides are in proximity, we assume that the variations on the input/drop waveguide and those on the ring waveguide are the same in a single MRR.

Considering (4), we define the total resonant-wavelength shift in an MRR ($T_{\Delta\lambda_R}$) by distinguishing the contribution of each variation to the resonant-wavelength shift in the MRR. A generalization based on passive and active MRRs can be made using the following model:

$$T_{\Delta\lambda_R}(\lambda_R, w', t', h', R') = \frac{\partial\lambda_R}{\partial w}(\sigma_w) + \frac{\partial\lambda_R}{\partial t}(\sigma_t) + \frac{\partial\lambda_R}{\partial h}(\sigma_h) + \frac{\partial\lambda_R}{\partial R}(\sigma_R). \quad (9)$$

Here, $\frac{\partial\lambda_R}{\partial w}$, $\frac{\partial\lambda_R}{\partial t}$, $\frac{\partial\lambda_R}{\partial h}$, and $\frac{\partial\lambda_R}{\partial R}$ denote the rate of changes in the MRR resonant wavelength with respect to the variations in the waveguide width, SOI thickness, slab thickness, and radius, respectively (i.e., *resonant-wavelength shift slopes*). Note that $\frac{\partial\lambda_R}{\partial h} = 0$ in the case of passive MRRs using strip waveguides. Moreover, σ_w , σ_t , σ_h , and σ_R are the standard deviations associated with the variations in the waveguide width, SOI thickness, slab thickness, and radius, respectively. Our prior work in [9], [14], [32], and that from other groups [10], [23] showed that the impact of different variations in MRRs can be combined linearly (see (9)). Note that $\nu_{w,t,h,R}$ in (4) can be initialized based on these standard deviations. Moreover, such standard deviations can be quantified through various fabrications or obtained from a silicon photonic fabrication vendor. The resonant-wavelength shift slopes ($\frac{\partial\lambda_R}{\partial w}$, $\frac{\partial\lambda_R}{\partial t}$, $\frac{\partial\lambda_R}{\partial h}$, and $\frac{\partial\lambda_R}{\partial R}$) in (9) can be approximated using a linear model as we found that the resonant wavelength in MRRs changes almost linearly under different variations (see Fig. 2(a)). For instance, using (4), $\frac{\partial\lambda_R}{\partial w}$ can be approximated as:

$$\frac{\partial\lambda_R}{\partial w} \approx \left| \frac{\Delta\lambda_R(\lambda_{R0}, w + \epsilon, t, h, R) - \Delta\lambda_R(\lambda_{R0}, w - \epsilon, t, h, R)}{2\epsilon} \right|, \quad (10)$$

where ϵ is an arbitrary parameter denoting a slight change in the waveguide width. Similarly, $\frac{\partial\lambda_R}{\partial t}$, $\frac{\partial\lambda_R}{\partial h}$, and $\frac{\partial\lambda_R}{\partial R}$ can be approximated.

As shown in Fig. 3, the design space of MRRs can be divided into physical-level parameters (e.g., waveguide width, SOI thickness, and slab thickness) as well as device-level parameters (e.g., radius and gap), all of which can be affected under FPVs. Among these parameters, only the waveguide width, MRR radius, and the gap can be determined (and explored) at design-time as the SOI and slab thickness are limited by the host wafer and available etching depths in the fabrication process. Note that we do not consider the gap in an MRR in our design-space exploration in this paper, but we account for the variations in the gap through waveguide width variations in the MRR (i.e., width variations in the coupling region in an MRR that impact the gap). Moreover, the impact of radius variations in MRRs is minimal (see Fig. 2(a)). Therefore, in the rest of this section an effort is made to explore the impact of altering the waveguide width at design-time on the MRR performance under FPVs. In particular, we consider two scenarios where (see Fig. 3): A. the input/drop waveguide width (w_i) and ring waveguide width (w_r) are the same, and B. they are different, as discussed next.

TABLE I: Different parameters used in our simulations

Parameter	Value	Parameter	Value
SOI thickness (t)	220 nm	Slab thickness (h)	90 nm
Radius (R)	10 μm	Wavelength (λ or λ_{R0})	1550 nm
σ_w	4.9 nm	σ_t	1.5 nm
σ_h	1.5 nm	σ_R	0.8 nm

A. MRR performance analysis when $w_i = w_r$

Using (4), (9), nominal design parameters in Table I, and Lumerical MODE [18], which is a commercial eigenmode solver and simulator, we quantitatively simulate the rate of changes in the resonant wavelength w.r.t. different variations (i.e., $\frac{\partial\lambda_R}{\partial w}$, $\frac{\partial\lambda_R}{\partial t}$, $\frac{\partial\lambda_R}{\partial h}$, and $\frac{\partial\lambda_R}{\partial R}$) in passive and active MRRs. Results for the fundamental transverse-electric (TE) mode are shown in Figs. 4(a) and 4(e) for passive and active MRRs, respectively. We vary the input (w_i) and ring (w_r) waveguide width together from 300 to 1500 nm (selected as an example). Note that increasing the waveguide width in MRRs will contribute to higher-order mode excitation (i.e., multi-mode waveguides). While multi-mode waveguides are beneficial in some applications (e.g., mode-division multiplexing), we will discuss feasible solutions to alleviate higher-order mode excitation in Section V. Therefore, our design-space exploration in this section focuses mainly on understanding the impact of increasing the waveguide width on the rate of changes in the MRR resonant wavelength due to the different variations.

One can observe from Figs. 4(a) and 4(e) that the resonant-wavelength shift slopes corresponding to the different variations are all different. In particular, the impact of waveguide width and slab thickness variations ($\frac{\partial\lambda_R}{\partial w}$ and $\frac{\partial\lambda_R}{\partial h}$) decreases as the waveguide width increases; this denotes that MRR's resonance tolerance to the variations in waveguide width and slab thickness improves as the waveguide width increases. On the other hand, the impact of SOI thickness variations remains high and dominant as the waveguide width increases in Figs. 4(a) and 4(e) (i.e., increasing the waveguide width does not have much impact here). Another observation is that the impact of radius variations on the resonant-wavelength shift is small in both passive and active MRRs (e.g., ≈ 0.12 nm shift in the resonance with 1 nm change in the MRR radius [21]). Note that the results in Figs. 4(a) and 4(e) are independent of the MRR radius and gap. Employing (9) and the σ values in Table I—obtained by averaging different σ values in Table III in Section VI—we also show the total resonant-wavelength shift ($T_{\Delta\lambda_R}$) in passive and active MRRs in Figs. 4(a) and 4(e), respectively (see the second y-axis). We observe that as the waveguide width increases, $T_{\Delta\lambda_R}$ decreases, thereby the MRR tolerance to different FPVs increases.

Leveraging (5) and the parameters in Table I, we analyze the cross-over coupling (κ) in passive and active MRRs with $g = 100$ and 200 nm while increasing the waveguide width ($w_i = w_r$) from 300 to 1500 nm. Results, shown in Figs. 4(b) and 4(f), indicate that while increasing the waveguide width helps improve an MRR tolerance to FPVs (see Figs. 4(a) and 4(e)), the cross-over coupling exponentially decreases as the waveguide width increases in the MRR. Accordingly, the drop-port response in MRRs with wider waveguides will be

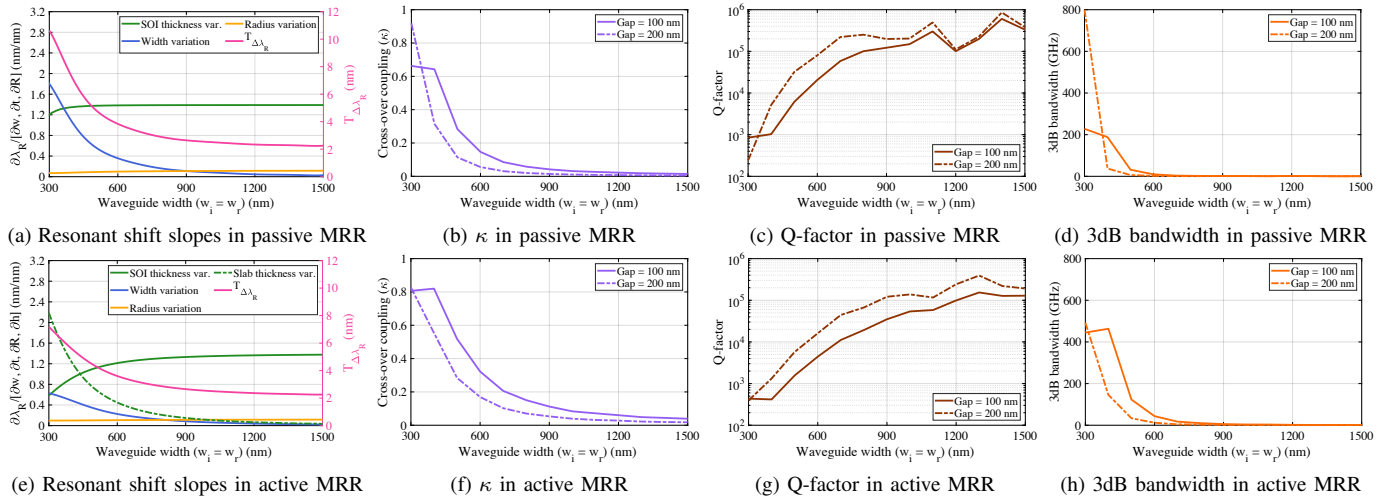


Fig. 4: Resonant-wavelength shift slopes and device performance in passive ((a)–(d)) and active ((e)–(h)) MRRs when the input/drop and ring waveguide widths increase from 300 to 1500 nm (x-axis). Here, (a) and (e) also show the total resonant-wavelength shift ($T_{\Delta\lambda_R}$). Results are for the fundamental TE mode with the parameters in Table I when $w_i = w_r$.

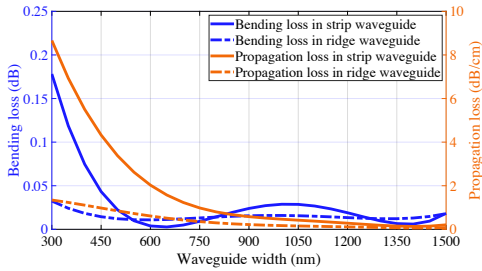


Fig. 5: Bending and propagation loss in strip and ridge waveguide as the waveguide width increases.

degraded. Note that the κ at $w_i = w_r \approx 300$ nm in Figs. 4(b) and 4(f) is higher for the gap of 200 nm as compared to that for the gap of 100 nm. We believe that this phenomenon is because of the optical signal being confined mostly in the coupling region rather than in the waveguide core for such an MRR design, similar to that of a slot waveguide [35].

As discussed in Section III, the cross-over coupling also determines the Q-factor and 3dB bandwidth in MRRs. Leveraging (6) and (8), Figs. 4(c), 4(d), 4(g), and 4(h) show the results for the Q-factor and 3dB bandwidth performance in passive and active MRRs when the MRR waveguide width ($w_i = w_r$) increases. As can be seen, as the MRR waveguide width increases, the Q-factor increases but the 3dB bandwidth considerably reduces to ≈ 0.1 GHz. Note that the Q-factor and 3dB bandwidth calculations in our work include the impact of increasing the waveguide width on the propagation loss and bending loss in MRRs. In particular, we estimate the waveguide propagation loss using the n_w model approximation [36], which is particularly useful if one only needs to quantify the relative comparisons of the propagation losses among different waveguide geometries under the same fabrication conditions [37]. Moreover, we estimated the changes in bending loss using Lumerical FDTD [18]. We found that both the propagation and bending losses decrease as the waveguide width increases, as shown in Fig. 5.

To summarize our findings in this sub-section, we demonstrate that one can improve the MRR tolerance to FPVs

by increasing the waveguide width in the MRR at design-time, but at the cost of severe reductions in the cross-over coupling that determines the drop-port response, Q-factor, and 3dB bandwidth in MRRs. To overcome this challenge while maintaining high MRR tolerance to FPVs, in the next sub-section an effort is made to study MRRs under FPVs while considering $w_i \neq w_r$ (i.e., an unconventional MRR).

B. MRR performance analysis when $w_i \neq w_r$

In this sub-section, we assume that the input and drop waveguide widths (i.e., w_i) are equal and can be different from the ring waveguide width (w_r)—see Fig. 3. The cross-over coupling in MRRs is proportional to the overlap and the interaction between the optical modes in the input and the ring waveguide. When increasing $w_i = w_r$ in an MRR, such an overlap reduces as the fundamental optical modes will be more confined inside the waveguide cores, and hence κ decreases. Therefore, here we examine a possible solution based on considering $w_i \neq w_r$ and then increasing w_r only (i.e., an *unconventional* MRR with $w_i < w_r$) to enhance κ in MRRs while improving the MRR tolerance to FPVs. Note that one can increase w_r adiabatically and similar to [27].

To enable design-space exploration in an unconventional MRR, we need to enhance the proposed analytical models in Section III to account for the effect of unequal input/drop and ring waveguide widths. In particular, the effective index in the MRR is the main parameter that needs to be recalculated when $w_i \neq w_r$ in an MRR. For simplicity and similar to the effective index analysis in waveguide grating structures [31], we consider the average effective index (n_{eff}) of the input/drop and the ring waveguide for our calculations in this sub-section, based on which all the analytical models proposed in Section III can be updated. Similar to Section IV-A, we also consider the impact of changes in bending and propagation losses in waveguides associated with unconventional MRRs (see Fig. 5) as the ring waveguide width increases.

Similar to Section IV-A and using the updated analytical models in Section III, Lumerical MODE [18], and the param-

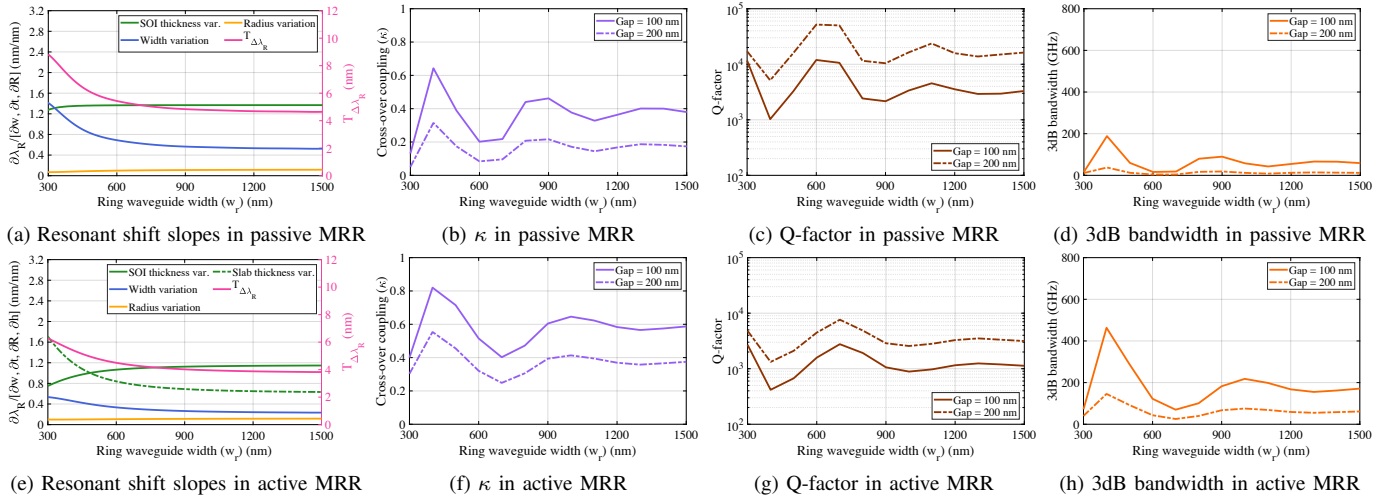


Fig. 6: Resonant-wavelength shift slopes and device performance in passive ((a)–(d)) and active ((e)–(h)) unconventional MRRs. Here, (a) and (e) also show the total resonant-wavelength shift ($T\Delta\lambda_R$). Results are for the fundamental TE mode with the parameters in Table I when $w_i = 400$ nm (considered as an example) and $w_i \neq w_r$. The x-axis shows the ring waveguide width (w_r) changes from 300 to 1500 nm.

eters in Table I, we explore MRRs with $w_i = 400$ nm (considered as an example) while increasing the ring waveguide width (w_r) from 300 nm to 1500 nm. Note that starting w_r from 300 nm is considered as an example and to show MRR performance when $w_r < w_i$. Figs. 6(a) and 6(e) show the resonant-wavelength shift slopes under the different variations in passive and active MRRs, respectively. As can be seen, the impacts of width, SOI thickness, slab thickness, and radius variations on the resonant wavelength of MRRs are similar to those shown in Figs. 4(a) and 4(e). However, compared to when $w_i = w_r$, the resonant-wavelength shift slopes, and hence the rate of changes in the resonant wavelength of an MRR under FPVs, are slightly higher when $w_i < w_r$. Consequently, the total resonant-wavelength shift—calculated using the σ parameters from Table I—in both passive and active MRRs is higher when $w_i < w_r$. Nevertheless, $T\Delta\lambda_R$ still decreases as w_r increases in Figs. 6(a) and 6(e), indicating that the proposed solution can help improve the tolerance of passive and active MRRs to FPVs.

Employing (5) and the parameters in Table I, Figs. 6(b) and 6(f) show the cross-over coupling in, respectively, passive and active MRRs for $g = 100$ and 200 nm when $w_i = 400$ nm and w_r increases from 300 to 1500 nm. Observe that κ increases at first when the ring waveguide width (w_r) increases from 300 nm to 400 nm, where the coupling is the highest at $w_i = w_r = 400$ nm. However, the coupling starts decreasing for $w_r > 400$ nm, until it increases again and reaches its second maximum where $w_r \approx 2w_i$. We simulated κ with other values for w_i and we found the same conclusion as the one in Figs. 6(b) and 6(f): κ increases when $w_r \approx \rho w_i$, where ρ is an integer. We further discuss such an increase in κ in Appendix A. Moreover, Figs. 6(c), 6(d), 6(g), and 6(h) show the Q-factor and 3dB bandwidth performance in unconventional passive and active MRRs while also including the changes in the propagation and bending losses in the MRRs (see our discussion in Section IV-A). A similar trend to κ can be observed in the Q-

factor and 3dB bandwidth. In particular, unlike when $w_i = w_r$ in Fig. 4, as w_r increases in an unconventional MRR, the decrease in Q-factor and 3dB bandwidth is inconstant. This is an interesting finding as it shows that an optimal MRR design is feasible with high tolerance to FPVs and also an acceptable cross-over coupling, Q-factor, and 3dB bandwidth (compare the results in Figs. 4 and 6).

To summarize our findings in this sub-section, we show that passive and active MRRs with unequal and variable w_i and w_r (called unconventional MRRs in this paper) can not only achieve high tolerance to FPVs but the cross-over coupling in these MRRs—and hence their drop-port response, Q-factor, and 3dB bandwidth—can be improved as well. Note that our conclusions and findings are independent of the assumption $w_i = 400$ nm, considered as an example in this sub-section, and we show results with other w_i values in Section V. To verify our results and findings in Sections IV-A and IV-B, we fabricate three different MRRs with variable w_i and w_r and evaluate the total resonant-wavelength shift and performance of our designs in the next sub-section.

C. Fabrication Results

We designed three passive TE-polarized MRR add-drop filters with a gap of 100 nm to experimentally validate our design-space exploration and results in this section. Fig. 7(a) shows the device layout of the different designed MRRs: MRR1 with $w_i = w_r = 400$ nm, MRR2 with $w_i = w_r = 800$ nm (see Fig. 4), and MRR3 with $w_i = 400$ nm and $w_r = 800$ nm (see Fig. 6). The radius (R) in MRR1 and MRR3 is set to be 10 μm . For MRR2, a racetrack resonator is considered with a radius and coupler length of 6 μm to ensure sufficient coupling between input/drop and ring waveguide (see κ in Fig. 4(b)). Note that the models developed in Section III can be easily extended to racetrack resonators (also see Appendix A). All the MRRs are designed to resonate at 1550 nm (i.e., desired/nominal resonance). Ten

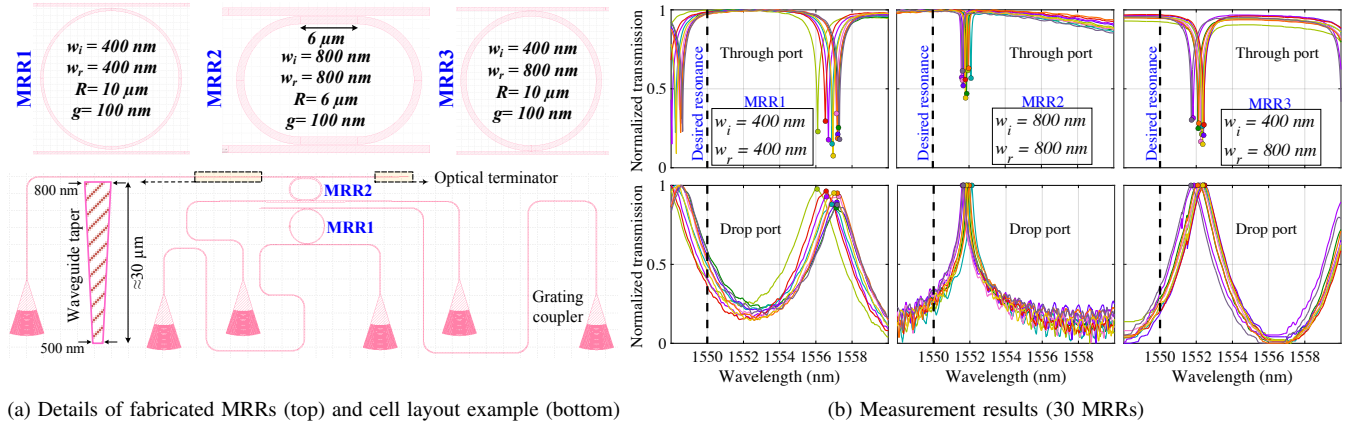


Fig. 7: (a) An example of the cell layout (bottom, MRR1 and MRR2 only) of the three designed passive TE-polarized MRRs with their design specifications (top). Note that 500-nm-wide strip waveguide is used for routing for all the MRRs. (b) Measured through- and drop-port responses obtained by testing 30 identical copies of MRRs in Fig. 7(a). All the MRRs were designed to resonate at 1550 nm, specified as desired resonance in the figures.

identical copies of each MRR were placed on a 1.5×0.6 mm² chip fabricated with a high-resolution electron-beam (EBeam) lithography system. Fig. 7(a) indicates an example of a unit cell of the designed MRRs with grating couplers designed for 1550 nm quasi-TE operation. Moreover, 220 nm thick SOI strip waveguides with a width of 500 nm are used for routing. Waveguide tapers of lengths 10–30 μ m are employed to ensure a single-mode operation.

Employing an automated silicon photonic testing station, we characterized the through- and drop-port responses of all the 30 MRRs (i.e., 10 identical copies per designed MRR), as shown in Fig. 7(b). Note that thermal variations were avoided using a thermal stage to preserve the chip temperature during the test. In Fig. 7(b), the desired resonant response is shown (at 1550 nm), and the resonant peaks that are specified by circles all belong to the same resonant mode. Table II summarizes the measurement results in MRR1–3. As can be seen, these experimental results are consistent with our design-space exploration results in Figs. 4 and 6. In particular, compared to MRR2 and MRR3, MRR1 has the highest cross-over coupling but lowest Q-factor of ≈ 500 (estimated using a Lorentzian fit) with an average total resonant-wavelength shift of 7.1 nm among the 10 resonant-wavelength peaks (i.e., the worst tolerance to FPVs among the fabricated MRRs). Moreover, MRR2 shows the best tolerance to FPVs as the average resonant-wavelength shift in MRR2 is only 1.8 nm. Nevertheless, MRR2 compromises cross-over coupling (see Fig. 4(b)), hence the noise on the drop-port response in Fig. 7(b), with the highest Q-factor of ≈ 1800 . MRR3, designed based on our proposed design-space exploration in Fig. 6, shows an average total resonant-wavelength shift of 2.1 nm (i.e., lower than MRR1 and slightly higher than MRR2), indicating a 70% increase in tolerance to FPVs (compared to MRR1) with a higher κ (compared to MRR2) and a Q-factor of ≈ 600 , which is higher than that for MRR1. Moreover, MRR3 has the best extinction ratio of 12 dB. It is worth mentioning that the difference between the calculated Q-factor results in Figs. 4 and 6 and those reported for our fabricated MRRs is due to the higher propagation (≈ 3 dB/cm) and coupler losses in our fabricated MRRs (note

TABLE II: Measurement results in MRR1-3

Parameter	MRR1	MRR2	MRR3
Total shift—Average	7.1 nm	1.8 nm	2.1 nm
Total shift—Standard deviation	0.38 nm	0.17 nm	0.25 nm
Q-factor—Lorentzian fit	500	1800	600

that our calculations and results in Figs. 4 and 6 consider a lossless coupler). Note that comparing our fabrication results with those from prior related work (e.g., [26], [27]) is unfair and invalid because the design parameters, objectives, and the fabrications are all different.

V. MRR DESIGN OPTIMIZATION UNDER FABRICATION-PROCESS VARIATIONS

Considering Figs. 4 and 6, employing wider waveguides in MRRs can help improve their tolerance to FPVs. However, considering different performance metrics altogether, several of which are contradicting, our design-space exploration in Section IV does not identify design parameters in MRRs (i.e., w_i and w_r in our paper) required to attain a high tolerance to FPVs while also achieving a desirable Q-factor and 3dB bandwidth and not imposing high overhead (e.g., silicon area) in MRRs. Indeed, choosing design parameters in MRRs based on a single performance metric may cause performance degradation when considering other metrics. For example, if an MRR is solely designed on the basis of having a high Q-factor, a designer might need to compensate in terms of photodiode sensitivity and also limited 3dB bandwidth of such a design. Leveraging the proposed MRR design-space exploration in Section IV, we present design optimization for passive and active MRRs under FPVs in this section.

The main goal in our optimization is to minimize the total resonant-wavelength shift in MRRs while considering MRR Q-factor and 3dB bandwidth. Note that κ , propagation loss, and bending loss are all considered in the optimization through their impact on Q-factor and 3dB bandwidth. Furthermore, to provide a designer with a peripheral vision while considering various properties of MRRs, we consider some additional metrics such as MRR footprint and higher-order mode excitation

(i.e., multi-mode waveguide effect) in MRRs. As a result, we can define a set of constraints (Ω) for the optimization as:

$$\Omega = \{(Q_m \leq Q \leq Q_M) \cap (B_m \leq B \leq B_M) \cap (A \leq A_M) \cap (\Delta n_{2nd} < \Delta n_{1st})\}. \quad (11)$$

Here, Q_m and B_m (Q_M and B_M) denote the minimum (maximum) acceptable Q-factor and 3dB bandwidth that can be determined by the designer. Note that $Q_{m,M}$ and $B_{m,M}$ are often application specific. Moreover, A_M is the maximum silicon-area overhead that is acceptable, and Δn_{1st} and Δn_{2nd} —considered to account for higher-order mode excitation—are the optical-confinement (guidance) strengths of, respectively, the first-order and the second-order TE mode in the MRR. Considering $\Delta n_{2nd} < \Delta n_{1st}$ ensures that the confinement strength of the fundamental mode is always stronger than that of the second-order mode (see our discussion below). We define the confinement (guidance) strength of an optical mode to be the difference between the effective index of that optical mode and the refractive index of cladding/substrate: $\Delta n = n_{eff} - 1.44$, where 1.44 is the refractive index of silicon dioxide at 1550 nm. The higher the confinement strength of an optical mode, the stronger the guidance of the optical mode will be in an MRR.

Algorithm 1 MRR Design Optimization under FPVs

```

1: Given  $R, g, \sigma_w, \sigma_t, \sigma_h$ , and  $\sigma_R$ 
2: for  $w_a \in [w_{i-\min}, w_{i-\max}]$  do /* Input waveguide */
3:   for  $w_b \in [w_{r-\min}, w_{r-\max}]$  do /* Ring waveguide */
4:      $Q_{w_a, w_b} \leftarrow$  calculate MRR Q-factor using (6)
5:      $B_{w_a, w_b} \leftarrow$  calculate MRR 3dB bandwidth using (8)
6:      $\Delta n_{2nd} \leftarrow \max(n_{eff_{2nd}}(w_a), n_{eff_{2nd}}(w_b)) - 1.44$ 
7:      $A_{w_a, w_b} \leftarrow$  calculate total silicon area in the MRR
8:      $C_{w_a, w_b} \leftarrow$  calculate  $T_{\Delta\lambda_R}$  with  $(w_a, w_b)$  using (9)
9:     if  $C_{w_a, w_b} < C^*$  ( $C_{w_a, w_b} \leq T_M$ ) then
10:       if  $\Omega_{w_a, w_b}$  is True then /* see (11) */
11:          $C^* \leftarrow C_{w_a, w_b}$ 
12:       Update  $s^*$  with  $(w_a, w_b)$  (add  $(w_a, w_b)$  to  $s$ )
```

For a given MRR design problem (i.e., R and g) and FPV parameters (i.e., $\sigma_w, \sigma_t, \sigma_h, \sigma_R$), an optimization search O can be used to find the optimal set of input and ring waveguide widths ($s^* = \{w_i^*, w_r^*\}$) with the minimum cost (C^*), where the cost function is the total resonant-wavelength shift in the MRR: $C = T_{\Delta\lambda_R}$, while satisfying the constraints in (11):

$$s^* \leftarrow O(s \in S, \text{Objective} : \min(C) \text{ s.t. } \Omega), \quad (12)$$

where S is a set of all the possible input and ring waveguide widths, defined by the designer. We use an exhaustive search approach, shown in Algorithm 1, to address the optimization in (12). Note that such an exhaustive search is made possible thanks to the high computational efficiency of the proposed models in Section III. To provide a designer with more flexibility when, for example, a total resonate-wavelength shift smaller than T_M is still acceptable (e.g., power budget allows for a tuning range up to T_M), the search can also find a set of input and ring waveguide widths (s) that satisfies the

constraints in (11) with $T_{\Delta\lambda_R} \leq T_M$. If this is desired, the parts underlined in Algorithm 1 must be considered. Note that we do not explore R and g in Algorithm 1 as the impact of R on $T_{\Delta\lambda_R}$ is negligible (see Fig. 2(a)) and g does not impact $T_{\Delta\lambda_R}$. Yet, leveraging the analytical models in Section III, the optimization search can be easily extended to include the impact of these parameters on the constraints defined in (11).

As an example for the MRR design optimization, we consider a passive (using strip waveguides) and an active (using ridge waveguides) MRR design problem with the design parameters in Table I and $g = 100$ and 200 nm for the passive and the active MRR, respectively. For brevity, we focus on the TE mode analysis. Similar to Section IV, we consider σ values listed in Table I. Note that all the parameters in this section are considered as an example, but the proposed method can be applied to any MRR design problem under FPVs. Leveraging the analytical models in Sections III and IV, we analyze the Q-factor, 3dB bandwidth, and total resonant-wavelength shift in the MRRs while sweeping both the input and ring waveguide widths from $w_{i-\min} = w_{r-\min} = 350$ nm to $w_{i-\max} = w_{r-\max} = 1200$ nm (see lines 2–3 in Algorithm 1). Moreover, we analyze the total silicon-area overhead and the second-order TE mode optical-confinement strength (see our discussion above) in the MRRs.

Figs. 8(a) and 9(a) show heatmaps for the Q-factor in the passive and active MRR, respectively. A low Q-factor in an MRR can increase the crosstalk noise and power penalty, and a very high Q-factor can put burden on the signal modulation (e.g., by enforcing the step size of the input signal to be smaller, which is limited by the tunable laser), and calls for a precise tuning mechanism in PICs [38]. As shown in Section III, MRR 3dB bandwidth and Q-factor are correlated. Assuming desired MRR 3dB bandwidth to be between 10 GHz and 50 GHz (see below), the MRR Q-factor should be larger than 3800 and smaller than 19000 (i.e., $Q_m = 3800$ and $Q_M = 19000$ in (11)). Accordingly, the design points (i.e., w_i and w_r) in Figs. 8(a) and 9(a) are specified using magenta squares. Figs. 8(b) and 9(b) show heatmaps for the 3dB bandwidth per wavelength in the passive and active MRR, respectively. A narrow 3dB bandwidth will result in heavy and undesired truncation of an optical signal spectrum, thus causing distortion, while a large 3dB bandwidth will result in higher crosstalk noise and power penalty. Similar to [7] and assuming a minimum signaling rate of 10 Gbps per optical channel (λ) for a wavelength-division multiplexing based link, we assume the desired bandwidth to be greater than 10 GHz and smaller than 50 GHz ($B_m = 10$ GHz and $B_M = 50$ GHz in (11)), based on which the design points in Figs. 8(b) and 9(b) are determined.

Figs. 8(c) and 9(c) show heatmaps for the silicon area consumption in the passive and active MRR, respectively. The silicon area in an MRR, which affects the fabrication cost, increases as w_i and w_r increase. Given the critical dimensions of an MRR add-drop filter, we define the silicon area as the sum of silicon-surface area on the input and drop waveguides and that on the ring. For simplicity, we consider the same silicon-area model for passive and active MRRs. As an example, we consider $w_i = w_r = 400$ nm—which

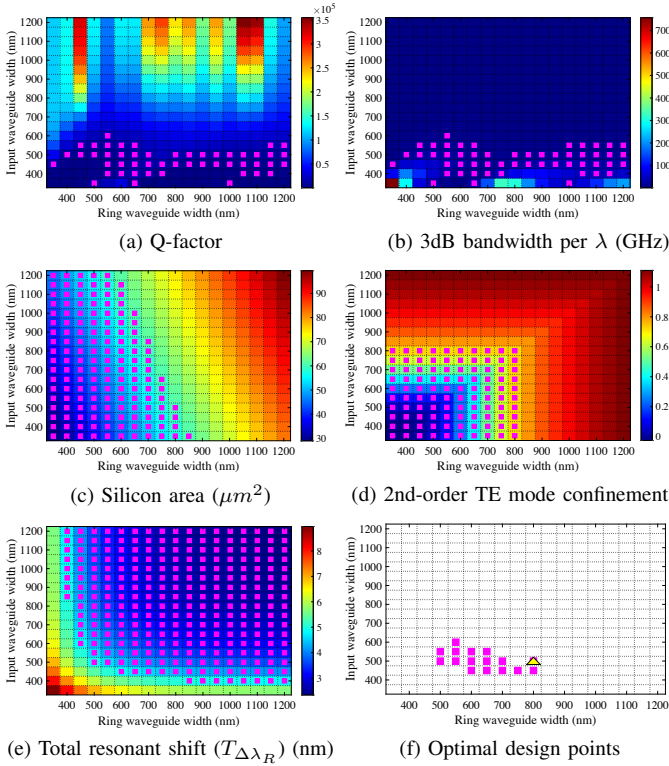


Fig. 8: Performance of a passive MRR designed using the parameters in Table I and with $g = 100$ nm where both the input and ring waveguide widths change from 350 to 1200 nm. The desired design points are selected and shown with magenta squares and the optimal MRR design region in (f) satisfies all the requirements in (a)–(e). Moreover, the yellow triangle in (f) shows the design point at which the total resonant-wavelength shift is minimum ($T_{\Delta\lambda_R} = 3.8$ nm).

corresponds to the total silicon area of $33.1 \mu\text{m}^2$ for the MRR case study considered in this section—as a baseline for silicon-area consumption. Accordingly, as w_i and w_r increase, we consider an upper limit for the resulting silicon area as twice the baseline (i.e., $A_M = 2 \times 33.1 \mu\text{m}^2$ in (11)), as indicated by the design points in Figs. 8(c) and 9(c).

Increasing the waveguide width will excite higher-order modes in MRRs (i.e., multi-mode waveguides). While this could be desired for applications such as mode-division multiplexing [39], PICs are often designed for single-mode operation. Note that higher-order mode excitation in our proposed MRR design can be avoided by engineering the MRR structure (e.g., adiabatically increasing the ring waveguide width [27]). Moreover, when increasing w_i and w_r , we ensure that the optical-confinement strength of the first-order TE mode (fundamental mode) is stronger than that of the second-order TE mode in both the passive and active MRRs. Figs. 8(d) and 9(d) show the optical-confinement strength of the second-order TE mode— Δn_{2nd} in (11) calculated based on line 6 in Algorithm 1—in the passive and active MRR, respectively. We consider 0.76 (for the passive MRR) and 0.99 (for the active MRR) as the optical-confinement strength of the first-order TE mode (Δn_{1st} in (11)) in the MRRs, calculated similar to line 6 in Algorithm 1 when $w_i = w_r = 400$ nm, the same baseline considered for the silicon area, and for n_{eff1st} . Note

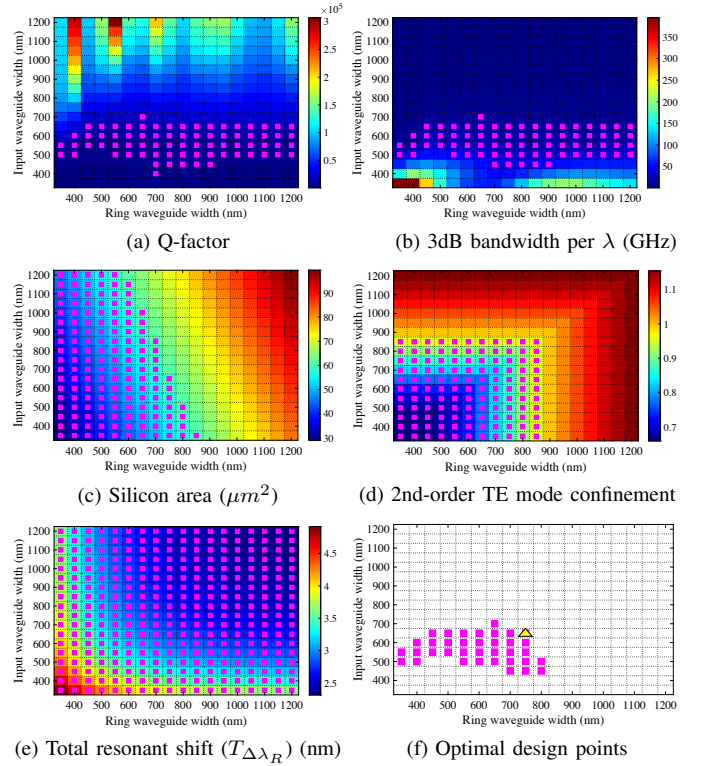


Fig. 9: Performance of an active MRR designed using the parameters in Table I and with $g = 200$ nm where both the input and ring waveguide widths change from 350 to 1200 nm. The desired design points are selected and shown with magenta squares and the optimal MRR design region in (f) satisfies all the requirements in (a)–(e). Moreover, the yellow triangle in (f) shows the design point at which the total resonant-wavelength shift is minimum ($T_{\Delta\lambda_R} = 2.8$ nm).

that Δn_{1st} increases as the waveguide width increases, hence Δn_{1st} at $w_i = w_r = 400$ nm (i.e., the baseline MRR) is considered in (11). The design points selected in Figs. 8(d) and 9(d) ensure that the optical-confinement strength of the second-order TE mode is always weaker than that of the first-order mode in the baseline (i.e., $\Delta n_{2nd} < 0.76$ and $\Delta n_{2nd} < 0.99$ for, respectively, the passive and active MRR in (11)).

Figs. 8(e) and 9(e) show heatmaps for the total resonant-wavelength shift in the passive and active MRR, respectively. As shown in Algorithm 1 and discussed above, a designer can consider a maximum tolerable total resonant-wavelength shift per MRR (T_M), and hence find a set of input and ring waveguide widths that corresponds to $T_{\Delta\lambda_R} \leq T_M$. Considering thermal tuning, state-of-the-art integrated heaters can consume as low as 27.5 mW/FSR [40] for resonance tuning in MRR add-drop filters. Assuming a maximum tuning power budget corresponding to FSR/2 nm per MRR (i.e., $T_M = \text{FSR}/2$ nm) and a similar tuning efficiency in passive and active MRRs, we select the design points in Figs. 8(e) and 9(e) where the tuning power consumption per MRR is less or equal to FSR/2 nm (i.e., 13.75 mW). This corresponds to the resonant shift of ≈ 5 and ≈ 6 nm for the passive and active MRR, respectively. Note that our assumption here (27.5 mW/FSR and $T_M = \text{FSR}/2$ nm) is considered as an example and it can be simply updated based on the tuning power budget in a system.

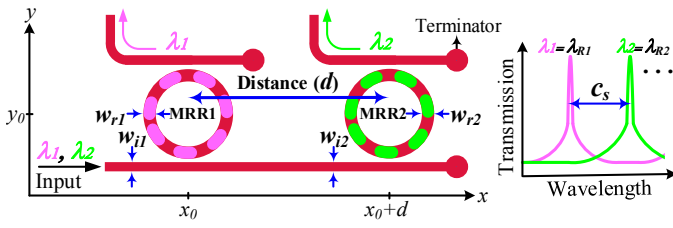


Fig. 10: A two-channel passive wavelength-selective MRR-based demultiplexer with two MRRs placed at a distance d and with a channel spacing c_s , designed based on Table I. Here, we assume $\lambda_{R1} = 1550$ nm and $\lambda_{R2} = 1553$ nm (radius in MRR2 is slightly different). Therefore, $c_s = 3$ nm.

TABLE III: Different parameters used to generate FPV maps

Design Parameter	Correlation Length	Standard Deviation
Waveguide width	$l_w = 4.5$ mm	Center: $\sigma_w = 4.2$ nm Edges: $\sigma_w = 5.5$ nm
SOI thickness	$l_t = 4.5$ mm	Center: $\sigma_t = 0.7$ nm Edges: $\sigma_t = 2.2$ nm
MRR radius	$l_R = 4.5$ mm	Center: $\sigma_R = 0.5$ nm Edges: $\sigma_R = 1$ nm

Figs. 8(f) and 9(f) show the search results in Algorithm 1 for the passive and active MRR, respectively. When T_M is considered (see above), the figures show the design points that satisfy the performance constraints discussed in this section for the considered examples of the passive and active MRRs. A designer can choose any design points (i.e., w_i and w_r) highlighted by magenta square in these figures to realize an MRR design with high tolerance to FPVs while achieving specific 3dB bandwidth and Q-factor, preserving silicon-area consumption, and alleviating higher-order mode excitation. When T_M is not considered, the search returns a single w_i and w_r per MRR (yellow triangles in Figs. 8(f) and 9(f)) that minimizes the total resonant-wavelength shift while satisfying all the constraints discussed in this section. Thanks to the low complexity of our proposed analytical models, the proposed design optimization can be easily integrated into an automated MRR design-space exploration and optimization tool.

VI. CASE STUDY: A WAVELENGTH-SELECTIVE MRR-BASED DEMULTIPLEXER

Leveraging the proposed MRR design-space exploration and optimization in Sections IV and V, here we improve the inter-device matching (i.e., channel-spacing accuracy) in a case study of a passive wavelength-selective MRR-based demultiplexer, shown in Fig. 10, under different FPVs. In addition, we develop virtual FPV wafer maps to account for actual layout information and fundamental variations in the waveguide width, SOI thickness, and radius, which are present on different length scales (i.e., correlations in variations).

We start by developing FPV wafer maps using a similar method proposed in [23] where an uncorrelated random distribution map with specific mean (μ) and standard deviation (σ) is first generated. Then, we convolve the resulting map with a Gaussian filter and specific correlation length (l) to obtain correlated FPV wafer maps. Moreover, we further enhance the correlated FPV maps by incorporating radial-variation effects:

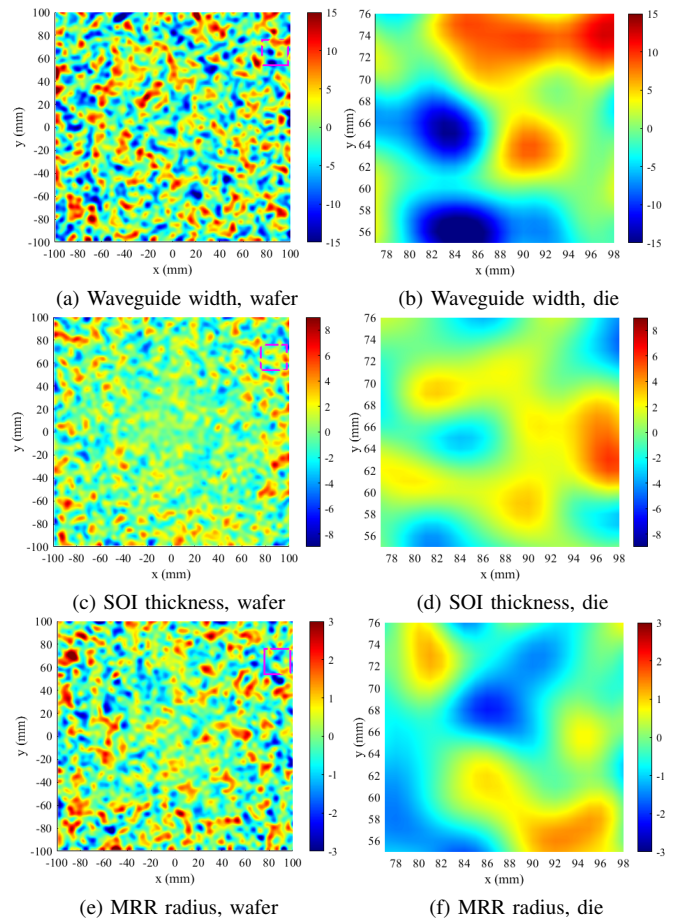


Fig. 11: Virtual FPV wafer maps ((a), (c), and (e)) and interpolated die maps ((b), (d), and (f)) that are correlated and mimic radial-variation effects. The maps are generated using the parameters in Table III with a mean of zero.

i.e., the non-uniformity increases as moving from the wafer center to the wafer edges, as reported also in other works [21], [27]; thus the wafer map center should have the least variations. To capture radial-variation effects, we first characterized the waveguide width and SOI thickness variations (i.e., σ_w and σ_t) at the center of several 200 mm wafers, and then repeated the same at multiple points while moving towards the edges of the wafers. The standard deviations were then averaged over the points within the same distance to the wafer center. For instance, Table III shows the standard deviations averaged at the center and edges of several 200 mm wafers that we characterized in collaboration with CEA-Leti (σ_w and σ_t only). Accordingly, we enhance our virtual wafer-map models with a variable standard deviation that increases almost linearly as moving from the center towards the wafer edges.

Leveraging the aforementioned method, we develop waveguide width, SOI thickness, and MRR radius virtual FPV maps with means of zero ($\mu_w = \mu_t = \mu_R = 0$) and correlation lengths (l_w , l_t , and l_R) and standard deviations (σ_w , σ_t , and σ_R) listed in Table III. In this table, $\sigma_{w,t}$ are analyzed through experimentally characterizing several 200 mm wafers at CEA-Leti, $l_{w,t}$ are taken from [23], and σ_R and l_R are considered as an example. Moreover, the table only shows the $\sigma_{t,w,R}$ at the wafer center and edges. Figs. 11(a), 11(c), and 11(e)

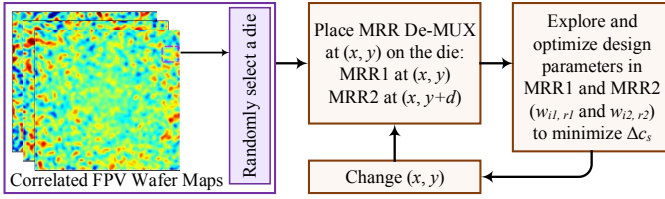


Fig. 12: An overview of the channel-spacing accuracy optimization for the MRR demultiplexer in Fig. 10.

show the resulting waveguide width, SOI thickness, and radius correlated wafer maps, respectively. Also, from each wafer map, we select a die with a size of $22 \times 22 \text{ mm}^2$, which is then interpolated as shown in Figs. 11(b), 11(d), and 11(f).

The channel spacing (c_s) in the two-channel MRR-based demultiplexer in Fig. 10 can be defined as the optical frequency space between the two MRRs' consecutive resonant wavelengths which can be given by $c_s = |\lambda_{R2} - \lambda_{R1}|$,

where $\lambda_{R2} > \lambda_{R1}$. As shown by Fig. 11, FPVs are present on different length scales and are correlated, hence we consider not only the MRR layout design parameters that are affected by different variations but also the positioning of each MRR on a die—FPVs on each MRR can be different—and the distance between the two MRRs (d in Fig. 10). Employing (9), we can model the deviated channel spacing (c'_s) in the two-channel MRR-based demultiplexer under different FPVs as:

$$c'_s = c_s + \Delta c_s, \quad (13a)$$

$$= c_s + |T_{\Delta\lambda_{R2}}(x + d, y) - T_{\Delta\lambda_{R1}}(x, y)|, \quad (13b)$$

where Δc_s denotes variations in the channel spacing. Also, $T_{\Delta\lambda_{Ri}}(x, y)$ is the total resonant-wavelength shift in MRR i located at position (x, y) on a die. We assume the FPVs on the input/drop and ring waveguides in the same MRR to be the same, but FPVs can be different in the two MRRs.

The channel-spacing accuracy under different FPVs can be improved by minimizing the channel-spacing variations (i.e., Δc_s in (13a)). Therefore, an optimization search similar to the one in (12) can be formed to design an MRR-based demultiplexer with high tolerance to FPVs, and hence high channel-spacing accuracy. Leveraging our proposed MRR design optimization discussed in Section V, this can be achieved by exploring and optimizing the input/drop and ring waveguide widths in MRR1 and MRR2 (i.e., $w_{i1, i2}$ and $w_{r1, r2}$ in Fig. 10) while considering the specific FPV profile experienced by each MRR. While we focus on the channel-spacing accuracy in this section, one can easily add other objectives (e.g., 3dB bandwidth and Q-factor) to the design-optimization problem, similar to the one proposed and addressed in Section V.

Considering (13b), the channel-spacing accuracy can be improved by applying $T_{\Delta\lambda_{R1}} \rightarrow 0$ and $T_{\Delta\lambda_{R2}} \rightarrow 0$, or $T_{\Delta\lambda_{R2/1}} \rightarrow T_{\Delta\lambda_{R1/2}}$, both minimizing Δc_s . Employing the FPV die maps in Fig. 11 ($22 \times 22 \text{ mm}^2$) and our MRR design optimization in Section V, we optimize the design parameters in MRR1 and MRR2—i.e., $w_{i1, r1}$ and $w_{i2, r2}$; other parameters are based on those in Table I—while uniformly positioning these MRRs at every location on the die, and analyzing channel-spacing variations (Δc_s) in the demultiplexer (see Fig. 12). We consider different scenarios where the MRRs

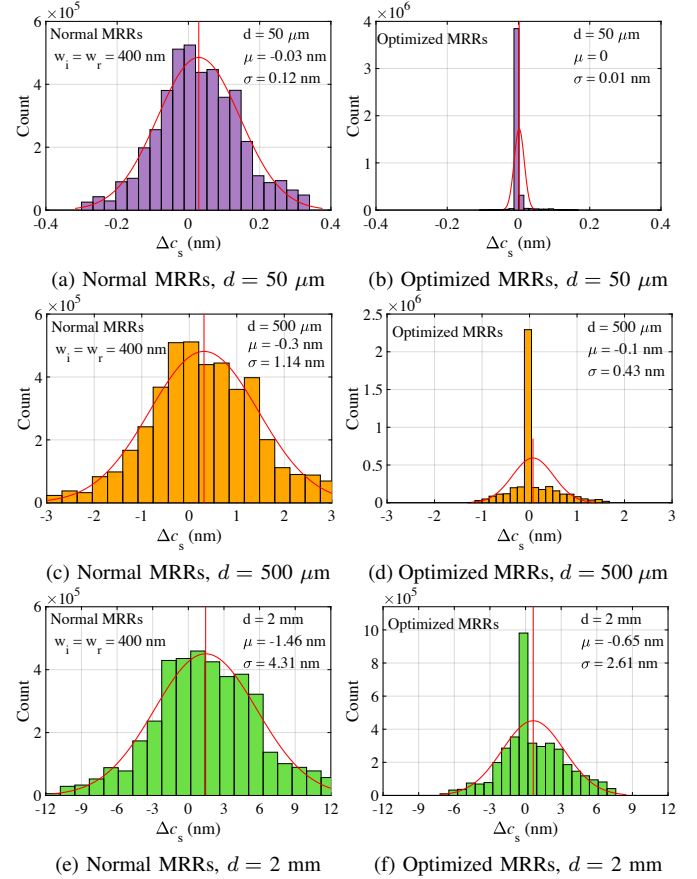


Fig. 13: Statistical analysis of channel-spacing variations (Δc_s) in the demultiplexer in Fig. 10 while considering normal and optimized MRRs and different distances (d) between the two MRRs. In the normal MRR design, $w_{i1, r1} = w_{i2, r2} = 400 \text{ nm}$. The optimized MRR design is based on the procedure in Fig. 12 with $w_{i1, i2} = 450 \text{ nm}$ and $w_{r1, r2} \in [570, 820] \text{ nm}$. The legends show the mean (μ) and standard deviation (σ) of the normal-distribution fit of each histogram. The nominal channel spacing in Fig. 10 is 3 nm.

are in proximity ($d = 50 \mu\text{m}$) and when they are placed apart on the die ($d = 500 \mu\text{m}$ and $d = 2 \text{ mm}$). Based on the optimal design region specified in Fig. 8(f), we set the input waveguides to be 450 nm wide in both MRR1 and MRR2 ($w_{i1, i2} = 450 \text{ nm}$), and explore ring waveguide width ($w_{r1, r2}$) from 570 to 820 nm to minimize channel-spacing variations. Therefore, the resulting MRR designs will satisfy all the performance requirements discussed in Fig. 8.

Fig. 13 indicates the channel-spacing variations (Δc_s) in the demultiplexer with normal (unoptimized) (Figs. 13(a), 13(c), and 13(e)) and optimized (Figs. 13(b), 13(d), and 13(f)) MRRs when $d = 50 \mu\text{m}$, $500 \mu\text{m}$, and 2 mm , and for 4×10^6 design samples. In the normal MRR design, $w_{i1, r1} = w_{i2, r2} = 400 \text{ nm}$. As can be seen, without MRR optimization, the channel-spacing variation is high and further increases as d increases (see Figs. 13(a), 13(c), and 13(e)). As shown in Figs. 13(b), 13(d), and 13(f), our MRR design optimization helps maintain the channel-spacing accuracy in $\approx 98\%$ of design samples within 0.05 nm when $d = 50 \mu\text{m}$, $\approx 80\%$ of design samples within 0.5 nm when $d = 500 \mu\text{m}$, and

$\approx 60\%$ of design samples within 2 nm when $d = 2$ nm. This clearly shows the effectiveness of our MRR design optimization. When $d = 50$ μm , both MRRs experience variations on a smaller length scale, hence intuitively inter-device matching is already higher and the optimization chooses MRRs of mostly equal width to improve channel-spacing accuracy. However, when d increases to 500 μm and 2 mm, MRRs experience much different variations and on a larger length scale, hence inter-device matching is more challenging. Nevertheless, the optimization succeeds to efficiently improve the channel-spacing accuracy even when d is large. Note that the optimization results in this section are in agreement with our design-space exploration results in Section IV. With such an optimized channel-spacing accuracy in MRR-based demultiplexers—enabled by our proposed MRR design optimization—one can compensate for wavelength shifts through collectively tuning all the MRRs, hence simplifying the circuit tuning and enhances its efficiency. The results presented in this paper show the promise of our proposed design-space exploration and optimization to improve MRR robustness in optical interconnects and emerging noncoherent artificial intelligence (AI) accelerators [41].

VII. CONCLUSION

With various advantages, silicon photonic microring resonators are often presented as the workhorse of emerging optical interconnects and photonic integrated circuits. In this paper, we present a comprehensive design-space exploration and optimization of MRRs under different fabrication-process variations. We consider variations in the waveguide width, SOI thickness, slab thickness (etching depth), and MRR radius in both passive and active MRRs. We present computationally efficient analytical models tailored to capture the impact of physical-level variations on MRR device-level performance. Leveraging these models, we exhaustively explore the design space of MRRs to enable an optimal MRR design with high tolerance to FPVs and desired Quality factor and 3dB bandwidth performance, all of which can be determined during design-time. As a case study, we apply our MRR design optimization to a two-channel wavelength-selective MRR-based demultiplexer where we show significant channel-spacing accuracy within 0.5 nm even when the MRRs are located 500 μm apart on a chip. Results in this work can help silicon photonic designers take into account the impact of FPVs during the design phase, thereby reducing the tuning power consumption and improving the circuit yield after fabrication. Also, our work shows the promise of design-space exploration in silicon photonic devices, paving the way for developing automated design tools and optimization techniques in this area.

APPENDIX A

CROSS-OVER COUPLING IN UNCONVENTIONAL MRRs

In Section IV, we show that the cross-over coupling (κ) in an MRR with wider waveguides and $w_i \neq w_r$ can be improved when $w_r \approx \rho w_i$, where ρ is an integer (see Figs. 3 and 6). Here, we analytically investigate such an improvement in κ by studying the cross-over coupling in a directional coupler (DC),

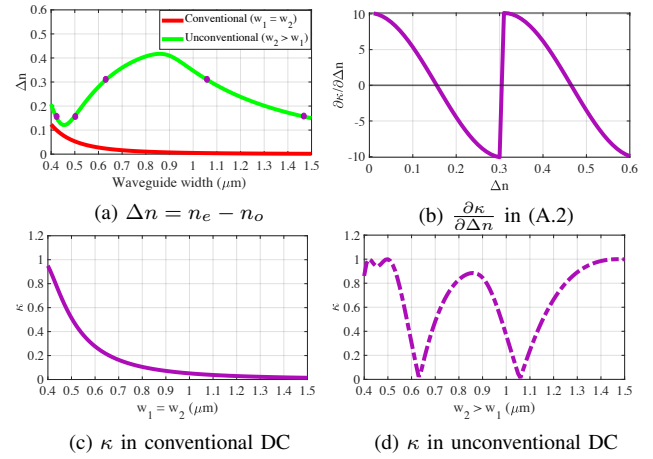


Fig. 14: (a) Difference between the effective indices of the symmetric (n_e) and antisymmetric (n_o) supermodes in a conventional and unconventional DC. (b) Rate of changes in the DC cross-over coupling (κ) w.r.t. the changes in the Δn . Cross-over coupling in (c) a conventional and (d) an unconventional DC (x-axis shows w_2). In these simulations, $L = 5$ μm and $g = 100$ nm. Also, $w_1 = 400$ nm in the unconventional DC.

which can represent the coupling region in an MRR. Moreover, we employ supermode theory [31] to analyze the cross-over coupling in the DC. Supermode analysis studies waveguides by considering the interfaces of the modes of the total structure. According to the supermode analysis, the effective indices of the first two eigenmodes of the coupled waveguides, which are known as symmetric and antisymmetric modes, determine the cross-over coupling in a DC. Considering supermode analysis and a DC with a gap of g and waveguide widths of w_1 and w_2 , the cross-over coupling can be defined as:

$$\kappa = \left| \sin \left(\frac{\pi \cdot \Delta n}{\lambda} \cdot L \right) \right|, \quad (\text{A.1})$$

where Δn is the difference between the effective indices of the symmetric (n_e) and antisymmetric (n_o) supermodes (i.e., $n_e - n_o$), λ is the wavelength, and L is the coupler length. Variations in the DC (e.g., increasing/decreasing the waveguide width) change Δn in (A.1). Considering the first derivative of the cross-over coupling (κ) with respect to Δn , we have:

$$\frac{\partial \kappa}{\partial \Delta n} = \left(\frac{\pi L}{2\lambda} \right) \frac{\sin \left(\frac{2\pi L}{\lambda} \cdot \Delta n \right)}{\left| \sin \left(\frac{\pi L}{\lambda} \cdot \Delta n \right) \right|}. \quad (\text{A.2})$$

As a case study, we quantitatively simulate κ and $\frac{\partial \kappa}{\partial \Delta n}$ in a DC with $L = 5$ μm and $g = 100$ nm, both considered as an example. We assume two scenarios: 1) a conventional DC in which $w_1 = w_2$, and 2) an unconventional DC in which $w_2 > w_1$. Using Lumerical MODE [18], Fig. 14(a) shows Δn where $w_1 = w_2$ increases from 400 to 1500 nm in the conventional DC, and w_2 increases over the same range in the unconventional DC with $w_1 = 400$ nm (considered as an example). Note that the gap is maintained at 100 nm. While increasing the waveguide width, Δn decreases when $w_1 = w_2$. However, when $w_2 > w_1$, Δn decreases at first and then it peaks when $w_2 \approx 2w_1$ in the unconventional DC in Fig. 14(a).

Employing (A.2), Fig. 14(b) shows $\frac{\partial \kappa}{\partial \Delta n}$ for Δn within the range $[0, 0.6]$ —see the y-axis in Fig. 14(a)—and with $L = 5 \mu\text{m}$. When $\frac{\partial \kappa}{\partial \Delta n} = 0$, we have a local maximum or minimum in κ . Figs. 14(c) and 14(d) show the cross-over coupling in the conventional and unconventional DC, respectively. As can be seen, when the waveguide width increases, κ decreases in the conventional DC, but it increases at some specific waveguide widths in the unconventional DC. In Fig. 14(a) and for the unconventional DC (when $w_2 > w_1$), we specified (see the circles in the figure) the waveguide widths corresponding to the Δn values at which $\frac{\partial \kappa}{\partial \Delta n} = 0$ in Fig. 14(b). Considering these waveguide widths (w_2), we can observe multiple maxima and minima in Fig. 14(d). In particular, $\frac{\partial^2 \kappa}{\partial \Delta n^2} = 0$ at $\Delta n = 0.3$ in Fig. 14(b) (i.e., the inflection point), corresponding to $w_2 = 630$ and 1050 nm (see Fig. 14(a) when $w_2 > w_1$) within which $w_2 \approx 2w_1$: there is a second maximum in κ in Fig. 14(d) when $630 \text{ nm} \leq w_2 \leq 1050 \text{ nm}$. Note that one can change L to show a similar trend and extend this to when $w_2 \approx \rho w_1$, where ρ is an integer. This explains the trend observed in the cross-over coupling in MRRs with $w_r > w_i$ (see Fig. 6). As discussed in Section V, such an increase in κ helps design MRRs which are not only tolerant to FPVs but also can achieve high Q-factor and 3dB bandwidth.

REFERENCES

- [1] J. E. Cunningham *et al.*, “Highly-efficient thermally-tuned resonant optical filters,” *Optics Express*, vol. 18, no. 18, pp. 19055–19063, 2010.
- [2] C. Manganelli *et al.*, “Large-FSR thermally tunable double-ring filters for WDM applications in silicon photonics,” *IEEE Photonics Journal*, vol. 9, no. 1, pp. 1–10, 2017.
- [3] T. Hu *et al.*, “Silicon photonic network-on-chip and enabling components,” *Science China Technological Sciences*, vol. 56, no. 3, pp. 543–553, 2013.
- [4] S. Pasricha and M. Nikdast, “A survey of silicon photonics for energy-efficient manycore computing,” *IEEE Design and Test*, vol. 37, no. 4, pp. 60–81, 2020.
- [5] W. Bogaerts *et al.*, “Silicon microring resonators,” *Laser and Photonics Reviews*, vol. 6, no. 1, pp. 47–73, 2012.
- [6] S. V. R. Chittamuru and S. Pasricha, “Improving crosstalk resilience with wavelength spacing in photonic crossbar-based network-on-chip architectures,” in *IEEE International Midwest Symposium on Circuits and Systems*, 2015, pp. 1–4.
- [7] M. Bahadori, M. Nikdast *et al.*, “Design space exploration of microring resonators in silicon photonic interconnects: Impact of the ring curvature,” *IEEE Journal of Lightwave Technology*, vol. 36, no. 13, pp. 2767–2782, Jul 2018.
- [8] S. K. Selvaraja *et al.*, “SOI thickness uniformity improvement using wafer-scale corrective etching for silicon nano-photonic device,” in *IEEE Photonics Society*, 2011, pp. 289–292.
- [9] M. Nikdast *et al.*, “Modeling fabrication non-uniformity in chip-scale silicon photonic interconnects,” in *IEEE/ACM Design, Automation and Test in Europe Conference and Exhibition*, 2016, pp. 115–120.
- [10] W. Bogaerts *et al.*, “Layout-aware variability analysis, yield prediction, and optimization in photonic integrated circuits,” *IEEE Journal of Selected Topics in Quantum Electronics*, vol. 25, no. 5, pp. 1–13, 2019.
- [11] N. M. H. Masri *et al.*, “WDM system based on radius variation of photonic microring resonators,” in *IEEE Student Conference on Research and Development*, 2017, pp. 243–246.
- [12] M. Bahadori *et al.*, “Thermal rectification of integrated microheaters for microring resonators in silicon photonics platform,” *IEEE Journal of Lightwave Technology*, vol. 36, no. 3, pp. 773–788, 2017.
- [13] S. Abel *et al.*, “A hybrid barium titanate–silicon photonics platform for ultraefficient electro-optic tuning,” *IEEE Journal of Lightwave Technology*, vol. 34, no. 8, pp. 1688–1693, 2016.
- [14] M. Nikdast *et al.*, “Chip-scale silicon photonic interconnects: A formal study on fabrication non-uniformity,” *IEEE Journal of Lightwave Technology*, vol. 34, no. 16, pp. 3682–3695, 2016.
- [15] F. Gan *et al.*, “Maximizing the thermo-optic tuning range of silicon photonic structures,” in *IEEE Photonics in Switching*, 2007, pp. 67–68.
- [16] W. A. Zortman *et al.*, “Silicon photonics manufacturing,” *Optics express*, vol. 18, no. 23, pp. 23 598–23 607, 2010.
- [17] H. Jayatilaka *et al.*, “Post-fabrication trimming of silicon photonic ring resonators at wafer-scale,” *Journal of Lightwave Technology*, vol. 39, no. 15, pp. 5083–5088, 2021.
- [18] Ansys Lumerical. [Online]. Available: <https://www.lumerical.com/products/>
- [19] Y. Xing *et al.*, “Hierarchical model for spatial variations of integrated photonics,” in *IEEE International Conference on Group IV Photonics (GFP)*, 2018, pp. 1–2.
- [20] R. G. Beausoleil *et al.*, “Devices and architectures for large-scale integrated silicon photonics circuits,” in *Optoelectronic Integrated Circuits XIII*, vol. 7942, 2011, p. 794204.
- [21] A. V. Krishnamoorthy *et al.*, “Exploiting CMOS manufacturing to reduce tuning requirements for resonant optical devices,” *IEEE Photonics Journal*, vol. 3, no. 3, pp. 567–579, 2011.
- [22] X. Chen *et al.*, “Process variation in silicon photonic devices,” *Applied optics*, vol. 52, no. 31, pp. 7638–7647, 2013.
- [23] Z. Lu *et al.*, “Performance prediction for silicon photonics integrated circuits with layout-dependent correlated manufacturing variability,” *Optics express*, vol. 25, no. 9, pp. 9712–9733, 2017.
- [24] Y. Wang *et al.*, “Characterization and applications of spatial variation models for silicon microring-based optical transceivers,” in *IEEE ACM/IEEE Design Automation Conference*, 2020, pp. 1–6.
- [25] Y. London *et al.*, “Behavioral model of silicon photonics microring with unequal ring and bus widths,” in *IEEE Optical Interconnects Conference*, 2019, pp. 1–2.
- [26] Y. Luo *et al.*, “A process-tolerant ring modulator based on multi-mode waveguides,” *IEEE Photonics Technology Letters*, vol. 28, no. 13, pp. 1391–1394, 2016.
- [27] Z. Su *et al.*, “Reduced wafer-scale frequency variation in adiabatic microring resonators,” in *IEEE/OSA Optical Fiber Communication Conference*, 2014, pp. 1–3.
- [28] A. Mirza, F. Sunny, S. Pasricha, and M. Nikdast, “Silicon photonic microring resonators: Design optimization under fabrication non-uniformity,” in *IEEE/ACM Design, Automation, and Test in Europe Conference and Exhibition*, 2020, pp. 484–489.
- [29] A. Mirza, S. Pasricha, and M. Nikdast, “Variation-aware inter-device matching in silicon photonic microring resonator demultiplexers,” in *IEEE Photonics Conference*, 2020, pp. 1–2.
- [30] G. Hocker and W. K. Burns, “Mode dispersion in diffused channel waveguides by the effective index method,” *Applied optics*, vol. 16, no. 1, pp. 113–118, 1977.
- [31] L. Chrostowski and M. Hochberg, *Silicon photonics design: from devices to systems*. Cambridge University Press, 2015.
- [32] M. Nikdast *et al.*, “Photonic integrated circuits: A study on process variations,” in *IEEE/OSA Optical Fiber Communication Conference*, 2016, pp. W2A–22.
- [33] M. Bahadori *et al.*, “Comprehensive design space exploration of silicon photonic interconnects,” *IEEE Journal of Lightwave Technology*, vol. 34, no. 12, pp. 2975–2987, 2016.
- [34] —, “Crosstalk penalty in microring-based silicon photonic interconnect systems,” *IEEE Journal of Lightwave Technology*, vol. 34, no. 17, pp. 4043–4052, 2016.
- [35] R. Hainberger, “Structural optimization of silicon-on-insulator slot waveguides,” *IEEE photonics technology letters*, vol. 18, no. 24, pp. 2557–2559, 2006.
- [36] D. Melati *et al.*, “A unified approach for radiative losses and backscattering in optical waveguides,” *Journal of Optics*, vol. 16, no. 5, p. 055502, 2014.
- [37] M. A. Tran *et al.*, “Ultra-low-loss silicon waveguides for heterogeneously integrated silicon/III-V photonics,” *Applied Sciences*, vol. 8, no. 7, 2018.
- [38] Y. Zhang *et al.*, “Design and demonstration of ultra-high-Q silicon microring resonator based on a multi-mode ridge waveguide,” *Optics letters*, vol. 43, no. 7, pp. 1586–1589, 2018.
- [39] B. A. Dorin and N. Y. Winnie, “Two-mode division multiplexing in a silicon-on-insulator ring resonator,” *Optics express*, vol. 22, no. 4, pp. 4547–4558, 2014.
- [40] P. Pintus *et al.*, “PWM-driven thermally tunable silicon microring resonators: Design, fabrication, and characterization,” *Laser & Photonics Reviews*, vol. 13, no. 9, p. 1800275, 2019.
- [41] F. Sunny, A. Mirza, M. Nikdast, and S. Pasricha, “Crosslight: A cross-layer optimized silicon photonic neural network accelerator,” *arXiv preprint arXiv:2102.06960*, 2021.