

Modélisation, caractérisation et contrôle d'opérations quantiques sur les qubits supraconducteurs

par

Élie Genois

Thèse présentée au département de physique
en vue de l'obtention du grade de docteur ès sciences (Ph.D.)

FACULTÉ des SCIENCES
UNIVERSITÉ de SHERBROOKE

Sherbrooke, Québec, Canada, 29 octobre 2024

Le Lundi, 28 octobre 2024

le jury a accepté la thèse de Élie Genois dans sa version finale.

Membres du jury

Professeur Alexandre Blais
Directeur de recherche
Département de physique

Professeure Eva Dupont-Ferrier
Membre interne
Département de physique

Professeur Emmanuel Flurin
Membre externe
Université Paris-Saclay, France

Professeur André-Marie Tremblay
Président rapporteur
Département de physique

Wisdom is nothing more profound than an
ability to follow one's own advice.
– *Sam Harris*

Sommaire

Les circuits supraconducteurs représentent une plateforme de choix pour le développement de technologies quantiques en raison de leurs propriétés de cohérence et de notre habilité à concevoir et à manipuler leur dynamique quantique de manière flexible. La réalisation de technologies ayant le potentiel de s'avérer utiles pour la société, telles qu'un ordinateur quantique tolérant aux fautes, nécessite cependant d'augmenter considérablement notre capacité à réaliser des opérations quantiques spécifiques sur ces systèmes, et ce sans engendrer d'erreurs. Pour ce faire, il est nécessaire d'améliorer à la fois notre caractérisation, notre modélisation et notre contrôle des dispositifs quantiques afin d'être en mesure de comprendre et d'atténuer de façon pratique les sources d'erreurs qui limitent les performances actuelles. Mes travaux s'attaquent à cette problématique en développant de nouvelles approches pour sonder les qubits supraconducteurs, pour utiliser l'information acquise afin de concevoir des descriptions simplifiées de leur fonctionnement et de leurs imperfections, puis finalement pour optimiser nos contrôles externes sur ces modèles afin de réaliser des opérations avec des erreurs minimales. Plus spécifiquement, je me suis d'abord intéressé à la modélisation et à la simulation de la première expérience de correction d'erreurs quantiques avec des qubits supraconducteurs, laquelle a été réalisée dans le groupe d'Andreas Wallraff à l'ETH Zürich (chapitre 3). Notre description numérique de cette expérience impliquant 17 qubits a permis de démontrer notre compréhension des résultats expérimentaux en plus d'étudier l'effet des différentes sources d'erreurs qui devront être mitigées dans les expériences futures. Dans l'objectif de contribuer concrètement à la réalisation d'opérations quantiques avec de meilleures fidélités, j'ai ensuite développé une méthodologie combinant l'apprentissage automatique et le contrôle optimal quantique (chapitre 4). Cette approche a le potentiel de s'avérer fort utile en raison de sa simplicité d'utilisation et de son efficacité, lesquelles sont démontrées ici sur le contrôle d'un qubit supraconducteur. En collaboration avec le groupe de recherche chez Google Quantum AI, je me suis ensuite intéressé à la caractérisation précise des opérations réalisées sur des dispositifs comportant des dizaines de qubits (chapitre 5). Nous avons développé deux protocoles

qui sont robustes et très efficaces afin d’extraire une description des portes logiques qui peut directement être utilisée afin de calibrer ces opérations. Finalement, j’ai contribué à la conception d’une approche numérique permettant de résoudre des problèmes arbitraires d’optimisation sur les dynamiques de grands systèmes quantiques ouverts (chapitre 6). Cette méthode possède un coût en mémoire minimal, ce qui nous permet de nous attaquer à des problèmes de contrôle optimal qui étaient autrement inaccessibles, comme l’optimisation de la mesure dispersive d’un qubit supraconducteur décrite par un modèle réaliste. Somme toute, mes travaux présentent plusieurs exemples où une approche basée sur un modèle et une compréhension physique nous permet de combiner la caractérisation au contrôle afin de contribuer concrètement au développement de technologies quantiques.

Mots-clés : Informatique quantique, Qubit supraconducteur, Simulation numérique, Caractérisation quantique, Apprentissage automatique, Contrôle optimal quantique

Remerciements

Je tiens à remercier profondément Alexandre pour sa présence d'esprit, sa sagesse, son souci du détail et sa gentillesse. Merci pour tous les efforts que tu déploies pour créer un environnement rempli de gens si talentueux et où c'est un réel plaisir de travailler. Merci de m'y avoir fait senti à ma place au cours des cinq dernières années. Je serai éternellement reconnaissant pour tout ce que tu as fait pour moi, et ce sera toujours un grand honneur pour moi de dire que j'ai fait mes études sous ta supervision.

Merci à tous les membres du groupe beaucoup plus intelligents que moi qui ont rendu ma vie au quotidien si agréable. Merci pour toutes les discussions scientifiques et totalement aléatoires, pour les fous rires, les jeux de société, les parties de volley et les bières au Siboire ou un peu partout ailleurs lors de conférences. Merci aux postdocs extraordinaires qui ont toujours été là pour me donner des conseils scientifiques et personnels : Cristóbal, Adrian, Alexander, Joachim, Alexandru, Manuel, Boris et Lev-Arcady, aux étudiantes et étudiants au doctorat avec qui j'ai eu la chance de partager cette expérience unique : Camille, Ross, Catherine, Lautaro, Othmane et Alex, aux étudiants et étudiantes à la maîtrise ou en visite qui ajoutent tellement de vie et de dynamisme au groupe : Simon, Marie-Frédérique, Benjamin, Ronan, Svend, Baptiste, Philippe et Adrien. Merci également à Benjamin pour ton leadership et ton savoir-faire et à Catherine pour le travail énorme que tu accomplis afin de faciliter nos vies.

I would especially like to thank Jonathan for his generosity, his patience, but most importantly for sharing with me his profound passion for science and rigorous understanding. Thank you for giving me the opportunity to work with such an amazing group of devoted and smart people at Google. You have been an amazing mentor and I am thrilled to get the chance of continuing to work with you. Thank you Dripto for welcoming me in your projects and in the whole LA experience, and for being such a good friend. Thank you Zhang for being a great manager, your perspectives and know-how are immensely useful in navigating everything that is going on in research.

Merci à l'amour de ma vie, Marilou, pour ton amour et ton appui inconditionnel à travers toutes ces épreuves. Merci d'avoir partagé tous ces beaux moments avec moi et de toujours m'encourager à donner le meilleur de moi-même. Merci à mes amis Olivier-Michel, Louis-Charles, Corinne, William, Karolann, Jean-Bastien, Dominic, Andy, Myriam, David, Allyson, Alexandra, Cedric, Sarah, Olivier, Antoine, Vincent et Kim pour toutes les belles fins de semaine, les voyages, les jeux, le spikeball, la bonne bière, le volleyball, le camping, les épluchettes et tous les autres plaisirs de la vie. Le fait de partager des moments avec vous me remplit toujours de bonheur et de gratitude d'être si bien entouré. Finalement, merci à ma famille adorée qui est toujours là pour moi et qui m'inspire profondément à devenir une personne généreuse qui fait le bien autour d'elle.

Table des matières

Sommaire	ii
1 Introduction	1
2 Les qubits supraconducteurs	5
2.1 Introduction	5
2.2 Les deux circuits étudiés : Le résonateur et le transmon	6
2.2.1 Le résonateur	6
2.2.2 Le transmon	8
2.3 Électrodynamique quantique en circuit : la mesure dispersive	11
2.4 Portes logiques sur le transmon	13
2.5 Simulations des dynamiques quantiques	18
3 Simulation précise de dynamiques quantiques à plusieurs transmons	21
3.1 La correction d’erreur quantique	22
3.2 Modèles de bruit en QEC	24
3.3 Méthodologie	26
3.3.1 Hamiltonien du système	27
3.3.2 Approche numérique	28
3.3.3 Opérations logiques	30
3.3.4 Modèle effectif 9+4	33
3.3.5 Modèle de bruit ajusté	36
3.4 Publication [1] : Réaliser la correction d’erreur quantique avec 17 transmons	37
3.5 Perspectives futures	47
4 Contrôle optimal basé sur l’apprentissage automatique	49
4.1 Mise en contexte	50
4.2 Concepts clés	53
4.2.1 Contrôle optimal d’un transmon	54

4.2.2	Apprentissage automatique	57
4.2.3	Apprentissage automatique pour le contrôle quantique	60
4.2.4	Contrôle optimal en deux étapes	62
4.3	Résultats principaux	64
4.4	Publication [2] : De la caractérisation par apprentissage automatique au contrôle optimal quantique	65
4.5	Perspectives	83
4.5.1	Modularité et compromis biais-variance	83
4.5.2	Opérations quantiques plus complexes	86
4.5.3	Applications directes	89
5	Protocoles de caractérisation quantique	92
5.1	Mise en contexte	93
5.1.1	Motivations pour CAFE	94
5.1.2	Motivations pour MEADD	97
5.2	Résultats principaux	102
5.3	Publication [3] : Caractérisation de la fidélité des opérations quantiques (CAFE)	104
5.4	Publication [4] : Estimation des paramètres unitaires de portes quantiques (MEADD)	122
5.5	Perspectives	140
6	Optimisation de la mesure en circuit QED	142
6.1	Mise en contexte	143
6.2	Résultats principaux	145
6.2.1	Notre méthodologie	145
6.2.2	Mesure dispersive optimale d'un transmon	148
6.2.3	Réinitialisation optimale d'un transmon	150
6.3	Publication [5] : Contrôle optimal de grands systèmes quantiques ouverts .	151
6.4	Perspectives	168
6.4.1	Librairie DYNAMIQS	168
6.4.2	Exemple d'une application directe de la méthode développée	169
	Conclusion	170
	Bibliographie	173

Table des figures

2.1	Résonateur	7
2.2	Transmon	9
2.3	Mesure dispersive du transmon	12
2.4	Contrôle du transmon	14
3.1	Simulation du bruit	25
3.2	Circuit quantique de correction d'erreur	31
3.3	Modèle effectif 9+4	34
3.4	Modèle de bruit ajusté	36
4.1	Vue d'ensemble du projet	53
4.2	Approche standard de contrôle optimal d'un qubit supraconducteur	55
4.3	Biais de modèle et contrôle optimal quantique par boucle ouverte	57
4.4	Contrôle par rétroaction	61
4.5	Approche proposée de contrôle quantique en deux étapes	63
4.6	Compromis entre biais et variance	84
4.7	Apprendre à partir de l'équation maîtresse	85
4.8	Portes parallèles à un qubit	87
4.9	Applications de l'approche MLQOC	90
5.1	Contexte d'une opération quantique	95
5.2	Fluctuations temporelles des paramètres unitaires	99
5.3	Découplage Dynamique	100
6.1	DYNAMIQS	170

Chapitre 1

Introduction

Née du désir de comprendre le fonctionnement du monde naturel qui nous entoure, la physique nous amène à concevoir de nouveaux outils nous permettant de sonder les phénomènes dont une description complète nous échappe. Une compréhension approfondie de ces phénomènes et des outils eux-mêmes ouvre la porte au développement de technologies innovantes ayant des applications bien au-delà des explorations scientifiques initialement poursuivies. Alors que cette compréhension est essentielle, c'est le *contrôle* qui permet concrètement de transformer nos connaissances scientifiques en technologies utiles [6]. Le développement de l'ordinateur quantique en est le parfait exemple.

Plus de 40 ans après la proposition concrète de contrôler les phénomènes quantiques afin de simuler le comportement d'autres systèmes quantiques complexes par Richard Feynman et ses contemporains [7], le domaine de l'informatique quantique se situe aujourd'hui à un point charnière. Bien que l'ensemble des applications d'un ordinateur quantique universel ne soit pas connu ni entièrement compris, la possibilité de simuler des phénomènes qui sont simplement hors de portée de tout ordinateur classique afin de concevoir, entre autres, de nouveaux matériaux, médicaments et processus chimiques a mené à un engouement économique d'envergure [8]. Alors que le domaine de l'informatique quantique et les réalisations expérimentales ont énormément progressé dans les dernières années, d'importantes contributions scientifiques sont encore nécessaires afin de réaliser les promesses associées à l'utilité de l'ordinateur quantique [9].

Un enjeu majeur dans le développement d'un ordinateur quantique utile possédant des milliers de qubits et réalisant des millions d'opérations est le compromis entre la cohérence de l'information quantique et notre capacité à précisément la contrôler de manière arbitraire. En effet, l'information que l'on utilise pour le calcul quantique est contenue dans l'état fragile

de notre système quantique, c'est-à-dire de nos qubits. Cet état est sensible à de nombreuses interactions avec son environnement pouvant mener à sa décohérence et donc à des erreurs de calcul. Afin de préserver l'information quantique, il est alors nécessaire d'isoler notre système de toute interaction externe. Cependant, ce couplage avec l'environnement externe est nécessaire pour contrôler nos qubits et effectuer du calcul quantique, par exemple en leur envoyant des signaux électromagnétiques. Somme toute, il est crucial de trouver un juste équilibre dans notre conception des dispositifs quantiques et de nos opérations afin de produire un minimum d'erreurs tout en réalisant les opérations logiques voulues avec une durée réaliste.

Bien que fortement réduite, l'exigence associée à la réalisation d'opérations possédant de très hautes fidélités persiste avec l'utilisation de la correction d'erreur quantique, un ensemble de techniques permettant de corriger les erreurs de calcul à mesure qu'elles surviennent [10, 11]. Pour les qubits supraconducteurs par exemple, l'une des plateformes de technologies quantiques les plus prometteuses [12, 13], le niveau de compréhension actuel de la fabrication et de l'opération de dispositifs contenant une centaine de qubits permet d'atteindre des fidélités supérieures à 99 % pour la mesure et les portes à deux qubits et supérieures à 99.9 % pour les portes à un qubit [14, 15]. Afin de réaliser un ordinateur quantique tolérant aux fautes, la fidélité et potentiellement la durée des opérations doivent cependant encore être améliorées, et ce dans un dispositif comportant au moins cent fois plus qubits [16, 17].

Dans l'objectif d'obtenir une maîtrise plus précise des différents processus ayant lieu dans le système quantique, il est nécessaire d'améliorer à la fois notre caractérisation, notre modélisation et notre contrôle des dispositifs quantiques. En effet, à mesure que les sources de bruit dominantes sont caractérisées, comprises et mitigées par la conception de meilleurs dispositifs et opérations, de nouvelles sources d'erreurs deviennent dominantes et doivent à leur tour être prises en compte afin de réduire les taux d'erreurs davantage. Dans cette thèse, je m'intéresse particulièrement à ce processus où le développement technologique et l'avancement de notre compréhension physique des phénomènes quantiques sont intimement reliés.

Mes travaux portent donc principalement sur les façons dont on peut sonder un système quantique afin d'extraire de l'information utile pour comprendre son fonctionnement (caractérisation), utiliser cette information afin de concevoir une description simplifiée de ses dynamiques (modélisation), puis finalement optimiser les façons dont on interagit avec le système afin de réaliser des opérations ayant une erreur minimale (contrôle). Les opérations d'intérêt ici seront la mesure de l'état des qubits ainsi que les portes logiques impliquant un et deux qubits. M'inspirant de la philosophie scientifique selon laquelle [18]

« tous les modèles sont erronés, mais certains sont utiles »

je cherche à construire des modèles spécifiques pour les qubits supraconducteurs dont la description des dynamiques réelles est suffisamment précise pour décrire l'expérience réalisée avec une précision acceptable et/ou pour permettre la conception de nouvelles opérations avec une fidélité cible. Appelant une telle représentation un *bon modèle*, une question centrale à tous les travaux de cette thèse est donc :

Quel est un bon modèle pour décrire les opérations des qubits supraconducteurs ?

Sans surprise, nous verrons qu'il n'existe pas de réponse absolue à cette question, mais plutôt un ensemble de solutions qui dépendent fortement des opérations spécifiques et des fidélités que l'on souhaite atteindre. Alors que les solutions offertes dans cette thèse sont vouées à être remplacées par de nouveaux et *meilleurs* modèles à mesure que le domaine progresse, les méthodologies élaborées ici pour répondre concrètement à cette question ont le potentiel de rester au cœur des développements technologiques futurs.

Au chapitre 2, j'introduis d'abord les concepts clés derrière la modélisation physique des opérations sur les qubits supraconducteurs, en plus de présenter les approches numériques permettant de simuler la dynamique de tels systèmes. Construisant sur ces idées, je présente au chapitre 3 l'approche de simulation que nous avons développée afin de décrire une expérience de correction d'erreurs quantiques où 17 qubits sont utilisés afin de former un qubit logique. En plus de démontrer notre compréhension des dynamiques complexes en jeux dans une telle expérience, ces simulations nous permettent d'étudier l'effet des différentes sources d'erreurs limitant la performance du qubit logique et ainsi de guider les efforts expérimentaux futurs.

Au chapitre 4, je développe une approche réaliste et facile d'utilisation ayant pour but d'améliorer la fidélité et la rapidité des opérations réalisées sur un dispositif supraconducteur. Plus spécifiquement, j'adopte une méthode d'apprentissage automatique afin de caractériser le système quantique directement à partir des données expérimentales qu'il produit. Par la suite, j'optimise les opérations quantiques sur ce modèle entraîné en utilisant les outils de la théorie du contrôle optimal quantique. Je démontre également l'utilité d'une telle approche à l'aide de simulations numériques détaillées d'un qubit supraconducteur, en plus de proposer de nombreuses applications concrètes et avenues de recherche associées à cette méthodologie.

Au chapitre 5, j'introduis deux protocoles de caractérisation quantique ayant pour but d'obtenir de manière efficace une description précise et détaillée de certaines portes logiques réalisées dans l'expérience. Le premier protocole, intitulé CAFE, sert à caractériser la fidélité d'une opération logique tout en obtenant un bilan des contributions cohérentes

et incohérentes à l'erreur totale mesurée afin de guider les actions à entreprendre pour réaliser de meilleures opérations. Le second protocole, intitulé MEADD, permet d'estimer les paramètres unitaires des portes à deux qubits avec une précision inégalée en éliminant l'influence d'une source de bruit importante, la fluctuation temporelle de la fréquence des qubits, laquelle limite les autres protocoles existants. Ces travaux ont été réalisés dans le cadre de mon stage avec l'équipe de recherche de Google Quantum AI.

Au chapitre 6, j'étudie finalement le contrôle optimal quantique de la mesure des qubits supraconducteurs. Pour ce faire, nous développons une approche numérique permettant de résoudre des problèmes d'optimisation arbitraires sur les dynamiques décrites par une équation maîtresse de Lindblad. Alors que notre méthode possède un coût en mémoire numérique minimal, elle permet de s'attaquer à l'optimisation de processus dans des systèmes quantiques ouverts beaucoup plus grands qu'auparavant. Nous appliquons cette approche à l'optimisation de la mesure d'un qubit supraconducteur et démontrons des gains potentiels importants par rapport aux protocoles de mesure utilisés actuellement. Étant donné l'ampleur des applications potentielles de notre approche numérique, j'ai contribué au développement d'une librairie disponible librement qui se nomme `dynamics`, laquelle j'introduirai brièvement à la fin de cette thèse.

Chapitre 2

Les qubits supraconducteurs

2.1 Introduction

Les circuits supraconducteurs représentent une plateforme de choix pour étudier de multiples phénomènes quantiques à une échelle macroscopique, profitant à la fois de la supraconductivité et d'une facilité à concevoir des interactions spécifiques entre la lumière et la matière. D'un côté, la supraconductivité fait en sorte que le courant électrique peut circuler sans aucune dissipation dans de tels circuits, ce qui permet de préserver l'information contenue dans ces signaux sur de longues périodes et ainsi faciliter l'observation de phénomènes quantiques cohérents. Les dynamiques de l'état supraconducteur, un condensat formé d'un nombre macroscopique de paires de Cooper et obéissant aux statistiques bosoniques, peuvent être décrites avec un seul degré de liberté [19, 20]. L'existence d'une telle description collective fait en sorte que les processus quantiques, typiquement associés aux phénomènes physiques s'appliquant à l'échelle atomique, peuvent être mesurés en pratique à l'échelle macroscopique en étudiant les dynamiques des excitations électromagnétiques d'un circuit supraconducteur [21]. D'un autre côté, les circuits supraconducteurs bénéficient d'une grande flexibilité dans la conception de leurs propriétés et de leurs couplages avec d'autres éléments supraconducteurs, lesquels sont par exemple utilisés pour la mesure et le contrôle précis du système quantique. En effet, il est possible d'utiliser des techniques de microfabrication bien établies et des appareils de contrôle de signaux micro-ondes afin de réaliser et d'opérer diverses expérimentations quantiques à base de circuits supraconducteurs.

En plus de représenter une plateforme idéale pour étudier des processus naturels fondamentaux tels que les interactions entre la lumière et la matière, le contrôle précis des degrés

de liberté de tels systèmes ouvre la porte à de nombreuses applications technologiques. Afin d'accomplir un tel contrôle en pratique, nous verrons que plusieurs ingrédients sont nécessaires. Dans les sections qui suivent, j'introduis les principaux éléments établissant les fondements sur lesquels ma recherche est construite, notamment le résonateur, le qubit transmon et la mesure dispersive combinant ces deux éléments, en plus de survoler la réalisation d'opérations logiques sur ces systèmes ainsi que leur modélisation numérique. Sans aucune prétention d'offrir une introduction complète et rigoureuse à la physique des circuits supraconducteurs, je me concentre ici sur les concepts qui seront utiles pour le reste de la thèse et j'invite la lectrice et le lecteur intéressés à se référer aux excellentes revues des Réfs. [13, 22, 23, 24, 12, 25].

2.2 Les deux circuits étudiés : Le résonateur et le transmon

Les circuits supraconducteurs sont typiquement fabriqués en imprimant un patron spécifique d'un métal normal comme l'aluminium ou le niobium sur un matériau diélectrique possédant de faibles pertes comme du saphir ou du silicium [26]. L'utilisation d'un tel système dans un réfrigérateur à dilution permet d'atteindre une température d'environ 10 mK et ainsi d'opérer dans la phase supraconductrice du métal et de minimiser le bruit associé à un environnement de température finie ($k_b T \ll \hbar \omega_0$) [22]. Alors que de nombreuses familles de circuits supraconducteurs existent et continuent d'être développées [25, 24], l'approche la plus répandue et sur laquelle portent mes travaux de recherche consiste à utiliser un circuit nommé *transmon* [27] et de le mesurer à l'aide d'un résonateur micro-onde imprimé à proximité [28].

2.2.1 Le résonateur

L'une des façons les plus simples de réaliser un oscillateur harmonique quantique avec un circuit supraconducteur consiste à imprimer un guide d'onde coplanaire en deux dimensions où un métal conducteur est entouré de chaque côté par deux autres lignes de métal formant le retour à la terre. Les conditions frontières (ouvertes ou fermées) de cette structure font en sorte que les modes électromagnétiques normaux qui y sont confinés sont quantifiés [13]. Grâce à l'ingénierie des matériaux et de la géométrie de ce guide d'onde, il est possible de concevoir avec précision l'inductance (L) et la capacité électrique (C) de ce circuit, de sorte à obtenir un résonateur avec une fréquence donnée entre 5 et 15 GHz. Un exemple de la géométrie typiquement utilisée pour ces résonateurs est présenté à la figure Fig. 2.1(a).

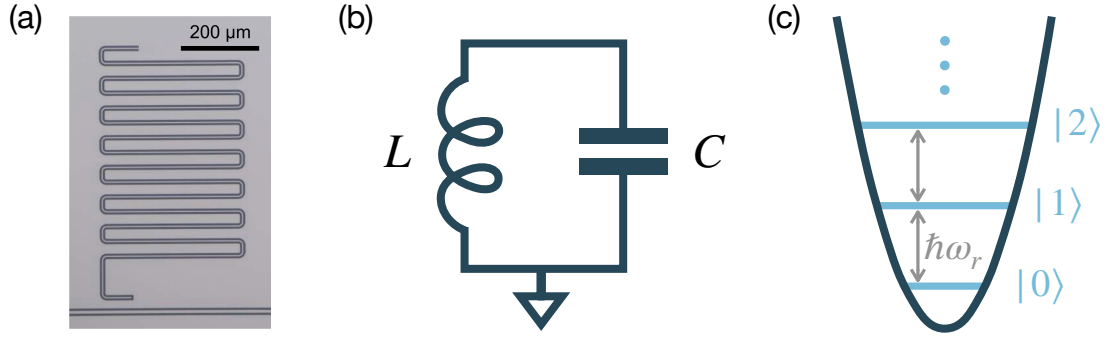


FIGURE 2.1 – Le résonateur. (a) Image optique tirée de la Réf. [26] d’un guide d’onde coplanaire (structure en forme de S) fait d’aluminium imprimé sur du silicium. Ce résonateur $\lambda/4$ est couplé capacitivement à une ligne de transmission (au bas de l’image) et possède une condition frontière ouverte à l’autre extrémité (haut de l’image). (b) Le modèle effectif du circuit illustré en (a) est celui d’un circuit LC, c’est-à-dire un simple oscillateur harmonique. (c) L’oscillateur harmonique quantique possède un puits de potentiel quadratique (bleu foncé) et des niveaux d’énergies discrets (bleu pâle) qui sont séparés par la même énergie $\hbar\omega_r$. Les panneaux (b) et (c) sont des adaptations d’une figure de la Réf. [13].

La description physique de ce circuit imprimé peut être obtenue à l’aide de la quantification d’un modèle classique, le modèle du télégrapheur, où le résonateur est représenté par une série de blocs fonctionnels où chaque unité infinitésimale de longueur possède une inductance et une capacité vers la terre [29]. Dans l’approximation où seul le mode fondamental du résonateur contribue aux dynamiques étudiées, cette description mène directement à l’Hamiltonien de l’oscillateur harmonique quantique [13]

$$\hat{H}_r = \hbar\omega_r \hat{a}^\dagger \hat{a}, \quad (2.1)$$

où ω_r est la fréquence fondamentale du résonateur et $\hat{a}^\dagger$ ($\hat{a}$) l’opérateur de création (d’annihilation) habituel. Comme cet Hamiltonien est équivalent à celui d’un circuit LC, le guide d’onde coplanaire est communément appelé un *résonateur* et est modélisé directement de cette façon, tel qu’illustré à la Fig. 2.1(b).

Dans l’objectif de contrôler l’état quantique du résonateur, il est possible de lui appliquer un signal micro-onde de fréquence ω_d et de phase ϕ_d , par exemple en utilisant la ligne de transmission couplée capacitivement qui est illustrée à la Fig. 2.1(a). Ce signal de contrôle dépendant du temps peut être décrit de la façon suivante [13]

$$\hat{H}_d(t)/\hbar = \epsilon(t)\hat{a}^\dagger + \epsilon^*(t)\hat{a}, \quad (2.2)$$

où $\epsilon(t) = A(t) \exp(-i(\omega_d t + \phi_d))$ est l’amplitude du signal qui est vue par le résonateur.

L’Hamiltonien $\hat{H}_d$ génère des états cohérents dans le résonateur, alors qu’il possède

directement la forme du générateur de l'opérateur de déplacement [30]

$$\hat{D}(\alpha) = e^{\alpha \hat{a}^\dagger - \alpha^* \hat{a}}, \quad (2.3)$$

lequel produit un état cohérent arbitraire lorsqu'il est appliqué sur l'état fondamental du résonateur,

$$|\alpha\rangle = \hat{D}(\alpha)|0\rangle = e^{-|\alpha|^2/2} \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} |n\rangle, \quad (2.4)$$

tel qu'obtenu directement à l'aide d'un développement en série de Taylor de l'exponentielle de $\hat{D}(\alpha)$.

Étant donné que les dynamiques de tels états cohérents sont décrites par des trajectoires classiques, c'est-à-dire en utilisant des distributions de probabilité de forme purement gaussienne pour décrire les états quantiques occupés, ils ne peuvent être utilisés pour réaliser du calcul quantique universel. Il s'avère alors nécessaire d'introduire une certaine non-linéarité dans le système afin de générer des états et des opérations quantiques intéressantes pour l'informatique quantique. En effet, pour concevoir et effectuer des opérations sur l'unité de calcul de l'informatique quantique, le qubit, il est nécessaire d'avoir un système où l'on peut définir deux états quantiques précis et effectuer des transformations unitaires arbitraires entre ces états. Nous avons donc besoin d'introduire le second circuit supraconducteur d'intérêt : le transmon.

2.2.2 Le transmon

L'élément clé à la base du succès des qubits supraconducteurs est la jonction Josephson [31], le seul élément de circuit non dissipatif connu qui est non linéaire. La jonction Josephson est fabriquée en séparant deux électrodes supraconductrices avec une mince barrière isolante, formant alors un lien que les paires de Cooper peuvent franchir par effet tunnel. Considérant que chacune de ces électrodes forme un îlot supraconducteur, la présence de la jonction fait en sorte que la différence entre le nombre de paires de Cooper sur chacun des îlots peut changer dans le temps. Cette dynamique cohérente est connue comme l'effet Josephson et est décrite simplement par [32]

$$\hat{H}_{JJ} = -\frac{E_J}{2} \sum_{n=-\infty}^{\infty} |n\rangle\langle n+1| + |n+1\rangle\langle n|, \quad (2.5)$$

où E_J est l'énergie Josephson et l'état $|n\rangle$ décrit la différence entre le nombre de paires de Cooper occupant chaque électrode. Pour obtenir la forme standard de l'Hamiltonien

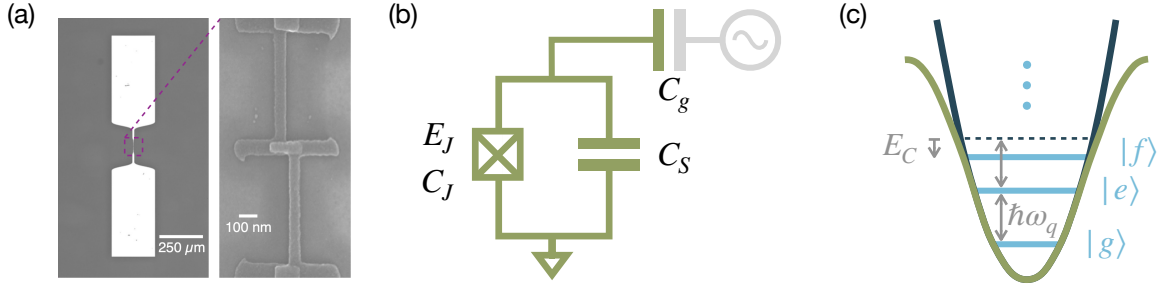


FIGURE 2.2 – Le transmon [27]. (a) Micrographie optique d’un transmon (gauche) et micrographie électronique à balayage de sa jonction Josephson (droite), laquelle est fabriquée en imprimant deux couches d’aluminium séparées par un diélectrique fait d’oxyde d’aluminium. Image tirée de la Réf. [33]. (b) Représentation sous forme de circuit du transmon. La jonction Josephson est décrite par son énergie E_J et sa capacitance intrinsèque C_J . La jonction est située en parallèle à une grande capacité C_S provenant des larges plaques de métal supraconducteur illustrées au panneau (a). (c) La jonction Josephson produit un potentiel périodique en forme de cosinus (vert), lequel est anharmonique et produit donc des états propres (bleu pâle) qui ne sont pas séparés par la même énergie. En particulier, l’énergie fondamentale $\hbar\omega_q$ de la transition $|g\rangle \leftrightarrow |e\rangle$ est environ E_C plus grande que l’énergie associée à la seconde transition $|e\rangle \leftrightarrow |f\rangle$, où E_C est l’énergie capacitive totale du circuit. Les panneaux (b) et (c) sont des adaptations de figures de la Réf. [13].

Josephson, on introduit la base des états de phase

$$|\varphi\rangle = \sum_{n=-\infty}^{\infty} e^{i\varphi n} |n\rangle, \quad (2.6)$$

où la différence de phase entre les électrodes est une valeur périodique $\varphi \in [-\pi, \pi)$. À partir de l’expression complémentaire

$$|n\rangle = \frac{1}{2\pi} \int_0^{2\pi} d\varphi e^{-i\varphi n} |\varphi\rangle, \quad (2.7)$$

on peut définir l’opérateur de phase

$$e^{i\hat{\varphi}} = \frac{1}{2\pi} \int_0^{2\pi} d\varphi e^{i\varphi} |\varphi\rangle \langle \varphi|, \quad (2.8)$$

lequel agit sur les états du nombre de paires de Cooper comme $e^{i\hat{\varphi}} |n\rangle = |n-1\rangle$. En substituant dans l’Éq. (2.5), on obtient alors directement l’Hamiltonien d’une jonction Josephson

$$\hat{H}_{JJ} = -E_J \cos(\hat{\varphi}). \quad (2.9)$$

Cet Hamiltonien peut être comparé directement à celui d’une inductance linéaire L , lequel prend la forme quadratique $\hat{H}_L = \hat{\Phi}^2 / 2L$, où $\hat{\Phi}$ est l’opérateur quantifié de flux ma-

gnétique traversant cette inductance. Comme illustré à la Fig. 2.2(c), le potentiel cosinusoidal de la jonction Josephson est anharmonique (c'est-à-dire non quadratique) et produit donc des états propres avec des énergies de transition différentes. De ce fait, les circuits supraconducteurs composés d'une ou de plusieurs jonctions Josephson possèdent des niveaux d'énergie que l'on peut directement adresser de manière distincte avec des contrôles de fréquence choisie. En raison de la ressemblance entre ce système et celui d'un atome avec des lignes spectrales distinctes, ce type de circuit supraconducteur est couramment appelé *atome artificiel* [34, 13].

L'atome artificiel le plus répandu est le transmon [27], un circuit où l'énergie capacitive est beaucoup plus faible que l'énergie Josephson, $E_J/E_C \gg 1$. Afin d'obtenir un circuit supraconducteur dans le régime du transmon, il suffit donc de combiner la jonction Josephson à grande capacitance en parallèle, par exemple en utilisant les deux grandes plaques de métal supraconducteur illustrées à la Fig. 2.2(a).

Le circuit électrique décrivant un tel système est présenté à la Fig. 2.2(b). Fixant l'un des îlots supraconducteurs du circuit à la terre, le nombre de paires de Cooper en surplus sur la seconde électrode peut être décrit par l'opérateur $\hat{n} = \hat{Q}/2e$, où $\hat{Q}$ est la charge électrique de l'îlot et e la charge d'un électron. Sachant que l'Hamiltonien associé à une capacité électrique C est donné par $\hat{H}_C = \hat{Q}^2/2C$ et utilisant $\hat{H}_{JJ}$ de l'Éq. (2.9), on obtient directement l'Hamiltonien du transmon

$$\hat{H}_T = 4E_C(\hat{n} - n_g)^2 - E_J \cos(\hat{\phi}), \quad (2.10)$$

où l'énergie capacitive $E_C = e^2/2C_\Sigma$ provient de la capacitance totale du circuit, soit $C_\Sigma = C_S + C_J + C_g$ et C_J est la capacitance propre de la jonction Josephson. La capacité C_g est utilisée pour coupler le transmon aux appareils électroniques de contrôle, par exemple à la source de tension illustrée en gris à la Fig. 2.2(b). Ce couplage permet de contrôler la charge de grille n_g du transmon et, comme nous le verrons à la section 2.4, de générer des opérations logiques sur celui-ci. Les opérateurs conjugués de différence de charge $\hat{n}$ et de phase $\hat{\phi}$ à travers la jonction obéissent à la relation de commutation attendue $[\hat{\phi}, \hat{n}] = i\hbar$.

Le principal avantage du transmon est que le régime $E_J/E_C \gg 1$ permet à ses états de basse énergie ($|g\rangle, |e\rangle, |f\rangle$) d'être exponentiellement insensibles au bruit de charge de grille, c'est-à-dire à des fluctuations temporelles de n_g dans le premier terme de l'Éq. (2.10) qui est beaucoup plus petit que le second terme [27]. Cette caractéristique est centrale au grand succès que connaît le transmon encore aujourd'hui, plus de 15 ans après son introduction [15, 14]. En effet, le bruit de charge de grille est omniprésent pour les circuits supraconducteurs et limite fortement la cohérence des qubits qui y sont sensibles, tels que

ceux ayant précédé le transmon [13]. Le transmon sera d'ailleurs au cœur de tous les travaux théoriques et réalisations expérimentales présentés dans cette thèse.

2.3 Électrodynamique quantique en circuit : la mesure dispersive

Ayant présenté le résonateur et le transmon individuellement, il est maintenant temps de les combiner afin de former l'architecture unie nommée électrodynamique quantique en circuit (cQED) [28, 35, 13]. Ce cadre théorique et expérimental est fondé sur la description et le contrôle des interactions fortes qu'il est possible de concevoir entre des qubits supraconducteurs, opérés par exemple en tant que systèmes à deux niveaux agissant comme des électrons (*matière*), et des photons micro-ondes provenant des champs électromagnétiques quantifiés qui sont confinés dans tels circuits (*lumière*), comme dans celui du résonateur. La physique décrivant un tel système est très riche et peut être utilisée afin de réaliser diverses applications en traitement de l'information quantique et bien au-delà de ce champ d'applications [13]. Dans le cadre de cette dissertation, nous allons étudier une seule application du couplage entre un qubit supraconducteur et un résonateur : la mesure dispersive du transmon.

Tel qu'illustré à la Fig. 2.3, l'idée pour mesurer l'état du transmon consiste à le coupler avec un mode auxiliaire dissipatif, le résonateur, où l'information s'échappe à un taux beaucoup plus grand que celui du qubit. Ce régime est réalisé en pratique en utilisant un grand couplage entre le résonateur et sa ligne de transmission. En raison de la forme du couplage entre le résonateur et le transmon, les deux états du qubit $|g\rangle$ et $|e\rangle$ ont des effets distinctifs sur le résonateur. Il est alors possible d'utiliser l'information provenant du résonateur afin de discriminer l'état du qubit. Voyons comment cette mesure fonctionne plus en détail.

Tout d'abord, le couplage capacitif entre le transmon et le résonateur crée un terme d'interaction entre l'opérateur de charge du qubit et celui du résonateur $\hat{n}_r \propto i(\hat{a}^\dagger - \hat{a})$, produisant ainsi l'Hamiltonien

$$\hat{H}_{tr} = 4E_C(\hat{n} - n_g)^2 - E_J \cos(\hat{\phi}) + \omega_r \hat{a}^\dagger \hat{a} - ig(\hat{n} - n_g)(\hat{a} - \hat{a}^\dagger), \quad (2.11)$$

où g représente la force du couplage capacitif entre le transmon et le résonateur [13]. Dans le cadre de notre étude de l'optimisation de la mesure dispersive au chapitre 6, nous utiliserons cet Hamiltonien complet afin de simuler les dynamiques temporelles du système. Nous ajouterons d'ailleurs un second résonateur en série agissant comme un filtre Purcell

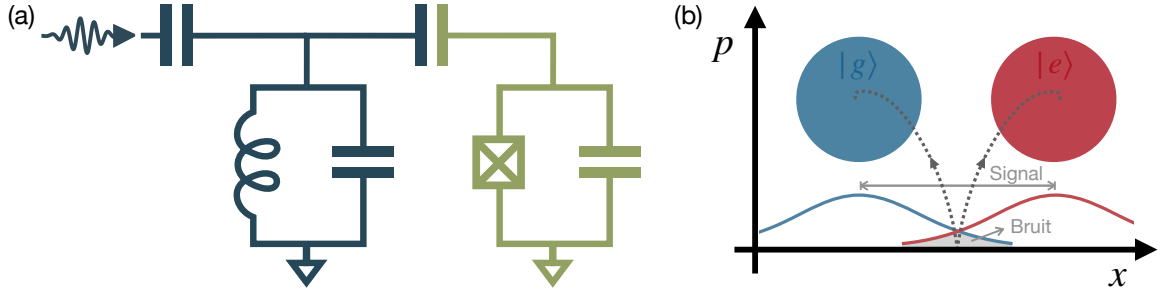


FIGURE 2.3 – Mesure dispersive du transmon à l’aide d’une approche d’électrodynamique quantique en circuit [13]. (a) En couplant capacitivement un résonateur dissipatif (bleu) et un transmon (vert), l’état du résonateur devient dépendant de l’état du qubit. Il est alors possible de sonder le résonateur avec un signal micro-onde et de discriminer les états $|g\rangle$ et $|e\rangle$ du transmon à partir du signal sortant du résonateur. (b) Dans l’espace de phase du résonateur, où x et p sont les coordonnées conjuguées, l’application du signal d’entrée produit des états cohérents différents dans le résonateur selon l’état du qubit. En principe, ces états cohérents peuvent être macroscopiques et donc mesurés directement en amplifiant le signal de sortie, par exemple à l’aide d’une détection homodyne (pas illustrée ici). Cette figure est une adaptation d’une figure de la Réf. [13].

permettant de préserver la cohérence du transmon [13]. Cet Hamiltonien est cependant complexe et il est difficile de voir directement comment le terme d’interaction ($\propto g$) nous permet de mesurer le qubit. Étudions donc une description simplifiée de ce système.

Afin d’utiliser le transmon comme un qubit ayant deux niveaux d’énergie qui sont cohérents et donc bien isolés de leur environnement, il est nécessaire d’éviter que ces états soient intriqués avec le mode dissipatif du résonateur lors des opérations logiques. Ainsi, on utilise typiquement un régime dit *dispersif* où le décalage de fréquence entre le transmon et le résonateur, $\Delta = \omega_q - \omega_r$, est beaucoup plus grand que l’amplitude de leur couplage g . Dans ce régime, il est possible de définir le petit paramètre $|\lambda| = |g/\Delta| \ll 1$ et d’utiliser une théorie de perturbations du second ordre afin d’obtenir un modèle grandement simplifié du système. En considérant le transmon comme un système à deux niveaux, une telle approche mène à l’Hamiltonien dispersif [28]

$$\hat{H}_{\text{disp}} \approx \frac{\hbar\omega'_q}{2}\hat{\sigma}_z + \hbar\omega'_r\hat{a}^\dagger\hat{a} + \hbar\chi\hat{a}^\dagger\hat{a}\hat{\sigma}_z, \quad (2.12)$$

où χ est le couplage dispersif et où les fréquences du qubit ω'_q et de la cavité ω'_r ont été renormalisées par l’interaction (avec des termes $\propto g^2/\Delta$) [27]. Le couplage dispersif s’exprime comme

$$\chi = -\frac{g^2 E_C / \hbar}{\Delta(\Delta - E_C / \hbar)}, \quad (2.13)$$

et possède typiquement une amplitude de l’ordre de quelques MHz.

En réexprimant l'Hamiltonien dispersif comme

$$\hat{H}_{\text{disp}} \approx \frac{\hbar\omega'_q}{2}\hat{\sigma}_z + \hbar(\omega'_r + \chi\hat{\sigma}_z)\hat{a}^\dagger\hat{a}, \quad (2.14)$$

on voit que la fréquence du résonateur (soit le préfacteur de $\hat{a}^\dagger\hat{a}$) dépend de l'état du qubit puisque la valeur attendue de l'opérateur $\hat{\sigma}_z$ est de ± 1 pour les états $|g\rangle$ et $|e\rangle$ du transmon. Tel qu'illustré à la Fig. 2.3, il est alors possible d'appliquer un signal de contrôle linéaire sur le résonateur afin d'y créer un état cohérent avec $\langle\hat{a}^\dagger\hat{a}\rangle = |\alpha|^2$ photons. Étant donné que le couplage dispersif $\propto \chi$ modifie la fréquence du résonateur, ces états cohérents vont tourner avec des sens de rotation différents dans le référentiel du résonateur (tournant à $\hbar\omega'_r$) de sorte à produire deux états cohérents distincts. Il est alors possible de mesurer l'état du résonateur, par exemple en utilisant une approche de détection homodyne [36], afin de discriminer si le transmon est dans son état $|g\rangle$ ou $|e\rangle$. J'invite la lectrice ou le lecteur intéressés à visiter les Réfs. [22, 37] pour plus de détails sur ce processus de mesure du résonateur.

Intuitivement, notre capacité à mesurer avec précision l'état du transmon, ou en d'autres mots de discriminer le bon état quantique avec une grande fidélité, est principalement relié au ratio entre le signal et le bruit (SNR). En adoptant un point de vue géométrique dans l'espace de phase du résonateur illustré à la Fig. 2.3(b), la mesure consiste à établir une frontière entre les états associés à $|g\rangle$ et $|e\rangle$ puis à étiqueter les états possibles de la cavité (tous les points dans l'espace de phase) comme étant associé à l'un ou l'autre de ces états. Dans cette description, les états cohérents correspondent à des distributions de probabilité en deux dimensions où le bruit de mesure est associé au volume de leur recouvrement alors que le signal produit est directement proportionnel à la distance entre le centre des deux états cohérents. C'est pourquoi au chapitre 6 nous chercherons à optimiser la séparabilité entre les états cohérents produits lorsque le transmon est dans $|g\rangle$ et $|e\rangle$, une mesure qui est quantifiée par le SNR.

2.4 Portes logiques sur le transmon

Ayant survolé la façon dont on peut définir un qubit dans les états de basse énergie d'un transmon et le mesurer avec précision en couplant ce circuit à un résonateur, il est maintenant temps de considérer les façons dont on peut effectuer des opérations logiques sur ces mêmes états quantiques. En effet, la réalisation de calculs quantiques arbitraires nécessite d'effectuer un ensemble universel de portes logiques, lequel est formé par exemple

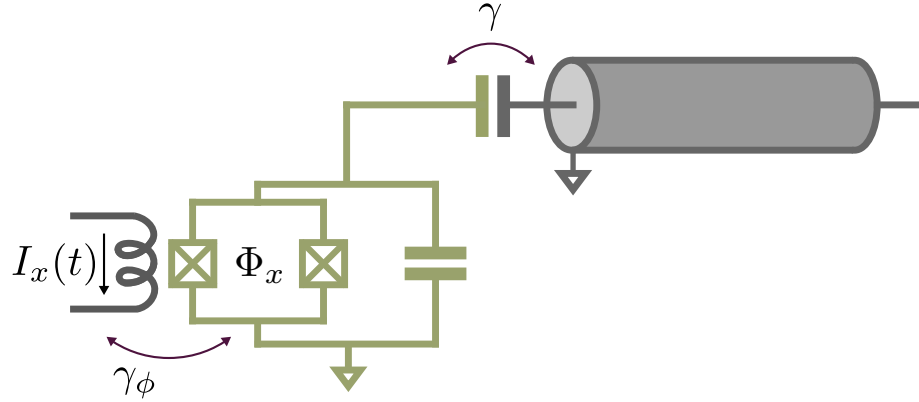


FIGURE 2.4 – Afin de contrôler les états quantiques du transmon, il est possible d'utiliser une ligne de transmission (gris) couplée capacitivement ou inductivement au circuit. L'application d'un signal de tension sur la ligne couplée capacitivement permet de modifier la charge de grille du transmon et de générer des dynamiques unitaires. De façon similaire, l'application d'un courant externe $I_x(t)$ permet de contrôler le flux traversant la boucle supraconductrice du transmon Φ_x et de modifier la fréquence et/ou l'état du transmon [27]. Le couplage du transmon aux composantes de contrôle l'expose également à un environnement externe, ce qui produit des canaux de dissipation illustrés avec le taux de relaxation γ et le taux de déphasage γ_ϕ . Figure tirée de la Réf. [13].

de rotations arbitraires à un qubit et d'un CNOT à deux qubits [38]. Toutes ces portes logiques sont représentées par des matrices unitaires de dimension $2^n \times 2^n$ agissant sur n qubits. En pratique, il est possible d'appliquer des signaux micro-ondes sur le transmon afin de réaliser de telles transformations unitaires. Spécifiquement, nous nous concentrons ici sur les opérations à un qubit et présenterons la porte à deux qubits utilisée et simulée dans cette thèse, le CZ, au chapitre 3.

Dans l'objectif de contrôler le transmon, il est possible d'envoyer des signaux électromagnétiques externes à travers des lignes de transmissions dédiées qui, selon la géométrie utilisée, peuvent se coupler capacitivement à l'opérateur de charge du transmon $\hat{n}$ ou inductivement à l'opérateur de phase $\hat{\phi}$. Alors que le couplage capacitif est naturellement réalisé en approchant la ligne de transmission à l'une des plaques du transmon, l'accès direct au degré de liberté de phase du transmon nécessite la création d'une boucle fermée de circuit supraconducteur. Tel qu'illustré à la Fig. 2.4, une telle boucle peut être réalisée en combinant deux jonctions Josephson en parallèle, un circuit communément appelé *SQUID* [39]. Comme nous le verrons aux chapitres 3 et 5, il est possible d'appliquer un flux externe constant à travers cette boucle afin de modifier la fréquence des transmons ou d'appliquer une modulation temporelle à ce flux afin de réaliser des opérations équivalentes à celles produites par un pilotage via l'opérateur de charge [27].

Étant donné la similitude entre le contrôle par la charge et la phase du transmon, on

se concentre ici sur la réalisation d'opérations en appliquant des signaux à une ligne de transmission couplée capacitivement au transmon. L'Hamiltonien de ce système piloté par un signal cohérent d'amplitude $A(t)$ est

$$\hat{H}_{1q}(t) = 4E_C(\hat{n} - n_g)^2 - E_J \cos(\hat{\phi}) + A(t) \cos(\omega_d t + \phi_d) \hat{n}, \quad (2.15)$$

où ω_d (ϕ_d) est la fréquence (phase) du signal.

Afin de comprendre comment l'Hamiltonien de l'Éq. (2.15) peut être utilisé pour réaliser des portes arbitraires sur le qubit, on utilise d'abord un développement en série de Taylor du terme $\cos(\hat{\phi})$ dans l'Hamiltonien du transmon profitant du fait que $E_J/E_C \gg 1$ et en négligeant l'effet de la charge de grille n_g comme on s'intéresse aux niveaux de basse énergie du transmon. On obtient alors

$$\hat{H}_q = 4E_C \hat{n}^2 + \frac{1}{2} E_J \hat{\phi}^2 - \frac{1}{4!} E_J \hat{\phi}^4. \quad (2.16)$$

En analogie directe avec l'oscillateur harmonique quantique, on introduit ensuite les opérateurs de création $\hat{b}^\dagger$ et d'annihilation $\hat{b}$ du transmon de sorte à obtenir

$$\hat{\phi} = \left(\frac{2E_C}{E_J} \right)^{1/4} (\hat{b}^\dagger + \hat{b}), \quad (2.17)$$

$$\hat{n} = \frac{i}{2} \left(\frac{E_J}{2E_C} \right)^{1/4} (\hat{b}^\dagger - \hat{b}), \quad (2.18)$$

où l'on voit que le régime transmon $E_J/E_C \gg 1$ justifie le développement en série pour $\hat{\phi}$. On obtient ainsi

$$\hat{H}_q = \sqrt{8E_C E_J} \hat{b}^\dagger \hat{b} - \frac{E_C}{12} (\hat{b}^\dagger + \hat{b})^4 \quad (2.19)$$

$$\approx \hbar \omega_q \hat{b}^\dagger \hat{b} - \frac{E_C}{2} \hat{b}^\dagger \hat{b}^\dagger \hat{b} \hat{b}, \quad (2.20)$$

où $\hbar \omega_q = \sqrt{8E_C E_J} - E_C$ et où nous avons fait appel à une approximation séculaire afin d'éliminer les termes où le nombre d'opérateurs $\hat{b}^\dagger$ est différent du nombre de $\hat{b}$. Cette approximation est justifiée par le fait que dans le référentiel tournant du qubit, obtenu avec à partir de la transformation unitaire $\hat{U}_q = e^{i\omega_q t \hat{b}^\dagger \hat{b}}$, ces termes oscillent avec une fréquence beaucoup plus rapide que leur amplitude, $\hbar \omega_q \gg E_C/4$, de sorte qu'ils se moyennent à zéro [13]. En exprimant l'Hamiltonien du contrôle, $A(t) \cos(\omega_d t + \phi_d)$, dans le référentiel

du qubit on obtient

$$\hat{H}_d(t) \approx \frac{i}{2} \left(\frac{2E_C}{E_J} \right)^{1/4} \frac{A(t)}{2} \left(\hat{b}^\dagger e^{i(\Delta t + \phi)} - \hat{b} e^{-i(\Delta t + \phi)} \right), \quad (2.21)$$

où $\Delta = \omega_d - \omega_q$ est le décalage de fréquence entre le signal de conduite et la transition $|g\rangle \leftrightarrow |e\rangle$ du transmon. Encore une fois, nous avons utilisé l'approximation séculaire afin d'éliminer les termes oscillants à la fréquence $\omega_d + \omega_q \gg |\tilde{A}(t)/2|$, où $\tilde{A}(t)$ inclut toutes les constantes (sauf un facteur 2 utilisé par convention) devant les opérateurs entre parenthèses.

Enfin, en effectuant une troncature aux deux premiers niveaux du transmon avec $\hat{b}^\dagger \rightarrow \hat{\sigma}_+$, $\hat{b} \rightarrow \hat{\sigma}_-$, et $\hat{\sigma}_\pm = (\hat{\sigma}_x \pm i\hat{\sigma}_y)/2$, on obtient

$$\hat{H}_d^{\text{qubit}}(t) = \frac{\tilde{A}(t)}{2} [\cos(\phi_d)\hat{\sigma}_x - \sin(\phi_d)\hat{\sigma}_y], \quad (2.22)$$

pour un signal résonant avec la fréquence du qubit ($\Delta = 0$). On comprend alors qu'en appliquant un signal de tension oscillant à la fréquence du qubit sur la ligne de transmission couplée au transmon, il est possible de réaliser des portes arbitraires en choisissant l'axe de rotation à l'aide de la phase ϕ_d ainsi que l'angle de rotation en contrôlant l'amplitude $A(t)$ et la durée du signal. Par exemple, on voit que l'on peut réaliser une porte

$$R_x(\theta) = e^{-i\theta\hat{\sigma}_x/2} = \cos(\theta/2)\mathbb{1} - i\sin(\theta/2)\hat{\sigma}_x, \quad (2.23)$$

en fixant la phase $\phi_d = 0$ et en appliquant l'amplitude constante $\tilde{A}(t) \equiv \Omega$ pour une durée T de sorte à obtenir l'opérateur d'évolution temporelle

$$\hat{U}(T) = e^{-i\hat{H}_d^{\text{qubit}}T/\hbar} = e^{-i\Omega T\hat{\sigma}_x/2}. \quad (2.24)$$

On peut alors directement identifier $\Omega T = \theta$. Ainsi, il est possible de réaliser une porte $R_x(\pi)$ en appliquant une amplitude effective de $\Omega/2\pi \approx 16.7$ MHz pendant 30 ns, des valeurs typiques pour les portes à un qubit réalisées avec des transmons.

Cependant, de nombreuses limitations nous empêchent de réaliser les opérations unitaires voulues avec une précision parfaite. En plus de la cohérence finie du transmon, de la largeur de bande limitée des impulsions qu'il est possible de produire en laboratoire et de l'impact des dynamiques que nous avons négligées pour arriver à l'Éq. (2.22), l'une des principales limitations à la qualité des portes logiques à un qubit provient de la fuite de niveau. En effet, le transmon ne constitue simplement pas un système à deux niveaux et l'opérateur de charge $\hat{n}$ possède des éléments de matrices significatifs couplant les états $|1\rangle \leftrightarrow |2\rangle$, $|2\rangle \leftrightarrow |3\rangle$, etc. Comme on peut comprendre à l'aide d'une transformée de Fourier,

les impulsions micro-ondes de courtes durées utilisées auront typiquement une contribution significative à la fréquence de la transition $|1\rangle \leftrightarrow |2\rangle$, laquelle est seulement décalée par 200 – 300 MHz de la transition $|0\rangle \leftrightarrow |1\rangle$ [13]. Cette contribution non négligeable produira un échange de population vers l'état $|2\rangle$ (et potentiellement les autres états) et cette fuite de niveau peut significativement réduire la fidélité des opérations réalisées [40].

Afin d'atténuer cet effet, il est cependant possible de concevoir soigneusement la forme des impulsions micro-ondes utilisées de sorte à minimiser la population sortant des deux premiers états. L'approche DRAG [41, 42] permet par exemple d'ajuster l'enveloppe des impulsions utilisées afin d'éliminer au premier ordre la contribution spectrale des contrôles à la fréquence de transition nuisible ($|1\rangle \leftrightarrow |2\rangle$). Cette approche simple s'avère très efficace pour atténuer la fuite de niveaux et est grandement répandue aujourd'hui [22]. Il ne s'agit là que d'une approche spécifique au contrôle de portes logiques et nous verrons au chapitre 4 que la théorie du contrôle optimal quantique fournit un ensemble d'outils puissants pour arriver à contrôler un dispositif quantique au meilleur de nos capacités.

Alors qu'il existe de nombreuses façons de quantifier la qualité avec laquelle un processus quantique $\mathcal{E}$ en approche un autre $\mathcal{U}$, lequel est typiquement une porte quantique (une opération unitaire) où $\mathcal{U}(\rho) = \mathcal{U}\rho\mathcal{U}^\dagger$, la quantité la plus souvent communiquée et celle que j'utiliserai dans les différents projets de cette thèse est la *fidélité moyenne de portes*, définie comme [43]

$$\mathcal{F}(\mathcal{E}, \mathcal{U}) = \int d\psi \langle \psi | \mathcal{U}^\dagger \mathcal{E}(|\psi\rangle\langle\psi|) \mathcal{U} | \psi \rangle, \quad (2.25)$$

où $0 \leq \mathcal{F} \leq 1$ et $\mathcal{E}$ est un canal quantique quelconque, c'est-à-dire une transformation complètement positive et préservant la trace de la matrice densité (CPTP). Cette métrique compare l'action du canal $\mathcal{E}$ à l'inverse de l'action du canal cible $\mathcal{U}$ sur *chaque* état pur de sorte à retrouver l'identité si les deux canaux sont équivalents pour cet état. La moyenne sur tous les états de l'espace d'Hilbert est ensuite utilisée, d'où la présence de l'intégrale. Ici, la mesure de Haar $d\psi$ est simplement une façon d'indiquer que tous les états que l'on échantillonne ont un poids égal et une probabilité normalisée telle que $\int d\psi = 1$ [44].

Évidemment, il n'est pas pratique de moyenner sur une infinité d'états purs $|\psi\rangle$. Heureusement, comme le canal $\mathcal{E}$ est une fonction quadratique de l'état $|\psi\rangle$, il s'avère qu'il suffit d'utiliser un ensemble d'états qui reproduit les statistiques de la mesure de Haar pour tout polynôme de degré deux afin de calculer la moyenne de façon exacte [45]. Un tel ensemble est appelé *2-design* et contient minimalement 4^n états purs pour n qubits. Comme nous le verrons au chapitre 5, cette propriété permet de grandement réduire le nombre d'expériences nécessaires afin d'estimer la fidélité d'une opération sur quelques qubits. En simulation,

comme ne nous sommes pas restreints à utiliser des matrices $|\psi\rangle\langle\psi|$ correspondant à des états physiques (c'est-à-dire des matrices hermitiennes de trace 1), il est possible de simplement utiliser un ensemble de matrices U_j formant une base orthonormale de l'espace des opérateurs unitaires, tel que les 4^n matrices de Pauli sur n qubits.

2.5 Simulations des dynamiques quantiques

Dans le cadre de cette thèse, on s'intéresse principalement à comprendre et à améliorer les opérations quantiques réalisées sur les qubits supraconducteurs. Pour ce faire, il est nécessaire d'avoir un modèle détaillé des caractéristiques du système, par exemple un Hamiltonien incluant les différentes interactions présentes dans le dispositif, puis d'étudier ses dynamiques temporelles sous l'application de nos contrôles externes. Étant donné la complexité du système et des contrôles changeant dans le temps, la description analytique des dynamiques est très limitée, voire impossible, et nous avons plutôt recours à des méthodes numériques.

L'approche numérique la plus simple consiste à utiliser les Hamiltoniens présentés jusqu'à maintenant, par exemple l'Éq. (2.15) pour les portes à un qubit et l'Éq. (2.11) pour la mesure, puis à résoudre l'équation de Schrödinger

$$i\hbar \frac{d}{dt} |\psi(t)\rangle = \hat{H}(t) |\psi(t)\rangle, \quad (2.26)$$

afin d'étudier l'impact des contrôles sur les dynamiques produites. Cette équation différentielle ordinaire (ÉDO) de premier ordre (linéaire et homogène) peut être résolue de plusieurs façons, par exemple en calculant le propagateur $\hat{U}(t)$ agissant comme $|\psi(t)\rangle = \hat{U}(t) |\psi(0)\rangle$, où

$$\hat{U}(t) = \mathcal{T} \exp \left(-i \int_0^t \hat{H}(t') dt' \right) \quad (2.27)$$

et $\mathcal{T}$ est l'opérateur de produit chronologique. Cette approche requiert typiquement une complexité numérique $\mathcal{O}(d^2)$ en mémoire pour entreposer l'Hamiltonien de dimension $d = 2^n$ pour n qubits et $\mathcal{O}(d^3)$ en temps pour calculer l'exponentielle d'une matrice. Une seconde approche consiste à profiter du fait que l'équation de Schrödinger est une ÉDO pour la résoudre en utilisant l'un des nombreux algorithmes d'intégration d'ÉDO existants et disponibles librement, comme une approche de Runge-Kutta [46]. Alors que la complexité en mémoire est la même, la complexité numérique en temps de cette approche est de $\mathcal{O}(d^2 \times m)$ où m est le nombre de pas de temps nécessaire, ce qui peut être avantageux dans certaines situations.

Cependant, pour décrire précisément le comportement observé en laboratoire, il est nécessaire de prendre en compte les nombreuses sources de bruits menant au temps de cohérence fini des circuits supraconducteurs. De plus, la dissipation peut jouer un rôle clé dans les dynamiques que l'on souhaite décrire, par exemple pour la mesure dispersive du transmon alors que la mesure d'un système quantique est un processus fondamentalement dissipatif [36]. On utilise une équation maîtresse afin de simuler de telles dynamiques quantiques, une méthode établie en optique quantique [30, 47]. Je réfère la lectrice et le lecteur intéressés à la Réf. [48] pour une introduction complète à la formulation d'une telle équation à partir des interactions entre un système et son environnement. Je me contente ici de résumer le fait que l'équation maîtresse de Lindblad que nous utilisons suppose que le système est couplé faiblement aux degrés de liberté de son environnement (approximation de Born) et que les dynamiques associées à ces derniers s'atténuent beaucoup plus rapidement que l'échelle de temps des dynamiques du système de sorte que l'information qui s'y échappe ne revient pas dans le système (approximation de Markov).

L'équation maîtresse de Lindblad décrivant l'évolution de l'état du système ouvert $\hat{\rho}$ prend la forme

$$\begin{aligned} \frac{d}{dt}\hat{\rho} &= -\frac{i}{\hbar} [\hat{H}, \hat{\rho}] + \sum_k \hat{L}_k \hat{\rho} \hat{L}_k^\dagger - \frac{1}{2} \left\{ \hat{L}_k^\dagger \hat{L}_k, \hat{\rho} \right\} \\ &= -\frac{i}{\hbar} [\hat{H}, \hat{\rho}] + \sum_k \mathcal{D}[\hat{L}_k] \hat{\rho} \\ &\equiv \mathcal{L} \hat{\rho}, \end{aligned} \tag{2.28}$$

où $\{\cdot, \cdot\}$ est l'anticommutateur, $\mathcal{D}[\cdot]$ le superopérateur de dissipation et $\mathcal{L}$ le Lindbladien. Les opérateurs de saut ou de dissipation $\hat{L}_k$ décrivent la manière dont l'environnement agit sur le système. Pour un circuit supraconducteur, ces opérateurs peuvent être obtenus à partir d'une description microscopique du système et de son environnement, par exemple en représentant l'environnement par une ligne de transmission de température finie qui est couplée capacitivement à un transmon. Une telle approche mène aux opérateurs de dissipation $L_k \in (\sqrt{\gamma(\bar{n}_\gamma + 1)}\hat{b}, \sqrt{\gamma\bar{n}_\gamma}\hat{b}^\dagger)$, où γ est le taux de relaxation du transmon dépendant de l'amplitude du couplage entre le circuit et l'environnement à la fréquence fondamentale du transmon et où $\bar{n}_\gamma$ est le nombre de photons thermiques présents dans l'environnement [13]. Une approche phénoménologique est également utilisée afin de décrire le déphasage que le qubit subit suite aux fluctuations temporelles des caractéristiques de son environnement. Cette description produit un troisième opérateur de dissipation de la forme $\sqrt{2\gamma_\varphi}\hat{b}^\dagger\hat{b}$, où γ_φ est le taux de déphasage pur du qubit [13]. Nous verrons plus en détail la forme des équations maîtresses utilisées dans chacun des chapitres qui suivent.

D'un point de vue fonctionnel, l'Éq. (2.28) constitue également une ÉDO du premier ordre qu'il est possible de résoudre numériquement à l'aide d'une grande variété d'approches. De façon similaire à la résolution de l'équation de Schrödinger, il est possible de calculer le propagateur qui est désormais un superopérateur agissant sur la matrice densité tel que

$$\hat{\rho}(t) = \mathcal{T} \exp \left(\int_0^t \mathcal{L}(t') dt' \right) \hat{\rho}(0). \quad (2.29)$$

Comme le Lindbladien $\mathcal{L}$ est représenté par une matrice de taille $d^2 \times d^2$, cette approche requière une complexité numérique $\mathcal{O}(d^4)$ en mémoire et $\mathcal{O}(d^6)$ en temps de calcul. Ces deux exigences deviennent rapidement problématiques pour la simulation de grands systèmes quantiques ouverts où $d = 2^n \gtrsim 500$. Dans le cadre de cette thèse, nous utilisons donc plutôt des algorithmes d'intégration d'ÉDO existant afin de résoudre l'équation maîtresse et d'obtenir les dynamiques quantiques de systèmes complexes. En évitant toute représentation sous forme de superopérateur, une telle approche nécessite une complexité $\mathcal{O}(d^2)$ en mémoire et $\mathcal{O}(d^3 \times m)$ en temps, où m est le nombre de pas d'intégration nécessaire. Au chapitre 3, nous utiliserons également une approche Monte-Carlo afin de réduire la complexité numérique en temps de calcul à $\mathcal{O}(d^2 \times m \times N)$, où N est une constante associée au nombre de trajectoires aléatoires qui sont échantillonnées.

Ayant survolé l'ensemble des concepts fondamentaux nécessaires au contrôle des qubits supraconducteurs et à la modélisation précise de telles dynamiques, il est maintenant temps d'explorer diverses applications de ces idées afin de faciliter la réalisation expérimentale d'un ordinateur quantique tolérant aux fautes.

Chapitre 3

Simulation précise de dynamiques quantiques à plusieurs transmons

Ce chapitre porte sur la simulation d’une expérience de correction d’erreur quantique impliquant 17 transmons réalisée par nos collaborateurs à l’ETH Zurich [1]. Cette expérience marque la toute première démonstration de correction d’erreur quantique avec des qubits supraconducteurs, quelques mois après une démonstration similaire dans un système de 10 ions piégés [49]. La simulation précise d’une expérience d’une telle envergure est évidemment très ardue, mais elle est également nécessaire afin de valider notre compréhension des résultats expérimentaux, de développer de nouvelles approches pour mieux contrôler un tel dispositif et de trouver les facteurs limitant la performance observée pour guider les efforts expérimentaux futurs. Mes travaux ont porté sur l’élaboration de ces simulations taillées sur mesure à l’expérience.

Après une courte introduction à la correction d’erreur quantique et au code de surface à la section 3.1, je présente les modèles de bruit typiquement utilisés pour simuler la performance d’un code quantique à la section 3.2. Par la suite, j’introduis notre méthodologie de simulation à la section 3.3 avant de présenter à la section 3.4 l’article scientifique qui partage et compare les résultats expérimentaux et de simulation. Je conclus en donnant quelques perspectives sur ce projet et sur les directions futures de recherche à la section 3.5.

3.1 La correction d'erreur quantique

Bien que de nombreuses applications intermédiaires de processeurs quantiques sont possibles, il est aujourd'hui largement reconnu que la correction d'erreur est nécessaire à la conception d'un ordinateur quantique tolérant aux fautes [9]. En l'absence de correction, il semble en effet improbable de concevoir un dispositif en mesure de réaliser des millions d'opérations sur des milliers de qubits avec des taux d'erreur sous le seuil de 10^{-10} , soit les ressources nécessaires estimées pour obtenir un avantage quantique afin de résoudre des problèmes intéressants [50, 51, 52, 53].

La correction d'erreur quantique utilise la redondance afin de protéger l'information quantique encodée dans plus de degrés de liberté qu'il ne serait nécessaire pour la représenter, en plus de réaliser du calcul quantique universel avec cette information. Au-delà de l'encodage de manière redondante et non locale dans un code de correction, il est nécessaire de s'assurer que l'information reste dans le bon état du code. Pour ce faire, une approche répandue consiste à mesurer des observables appelés des *stabilisateurs* [54] et à utiliser les syndromes classiques obtenus (les 0 et les 1 mesurés) afin de décoder l'état dans lequel le système se trouve ainsi que les opérations à appliquer afin de retourner dans un état du code, si nécessaire. Je note que de nombreuses approches alternatives existent, par exemple en concevant un système où la dissipation stabilise les états du code de manière autonome [55, 56]. On s'intéresse spécifiquement ici à l'opération d'un code de surface, le type de code stabilisateur le plus répandu [16].

De manière plus formelle, les états d'un code stabilisateur $|\psi_L\rangle$ appartiennent au sous-espace des états propres de valeur propre $+1$ d'un groupe abélien $\mathcal{S}$ appelé le groupe stabilisateur [54]

$$S_j|\psi_L\rangle = +1|\psi_L\rangle, \quad \forall S_j \in \mathcal{S}. \quad (3.1)$$

Pour un code dénoté comme $[[n, k, d]]$ utilisant n qubits physiques et encodant k qubits logiques avec une distance d , le groupe $\mathcal{S}$ est généré par $m = n - k$ opérateurs de Pauli indépendants. Ces stabilisateurs font partie du groupe des matrices du Pauli sur n qubits, $\mathcal{P}_n$. La distance d correspond au poids minimal d'un opérateur de Pauli qui commute avec tous les opérateurs de $\mathcal{S}$ sans être lui-même dans $\mathcal{S}$. Intuitivement, cet opérateur $E \in \mathcal{P}_n$ a un effet sur les états du code qui est non-trivial, mais qui n'est pas détectable alors qu'il ne crée que des syndromes triviaux

$$S_j(E|\psi_L\rangle) = ES_j|\psi_L\rangle = (+1)E|\psi_L\rangle, \quad \forall S_j \in \mathcal{S}. \quad (3.2)$$

Ainsi, le code quantique est en mesure de corriger toute erreur de Pauli ayant un poids

inférieur à d , où le poids correspond au nombre de qubits affectés de manière non triviale par l'opérateur. Afin d'implémenter un tel code stabilisateur, l'idée consiste à successivement mesurer l'ensemble des stabilisateurs du groupe $\mathcal{S}$, formant alors un *cycle* du protocole de correction, puis d'utiliser un décodeur, un algorithme classique, afin de déduire les erreurs qui se sont produites lors du cycle. Étant donné que la mesure d'un stabilisateur projette l'état logique dans le sous-espace de valeur propre -1 du stabilisateur S_j lorsqu'un syndrome non trivial est détecté, il suffit alors d'appliquer délibérément le bon opérateur E_j sur le système afin de corriger l'erreur qui est survenue étant donné que $E_j^2 = \mathbb{1}$ pour toute matrice de Pauli. La tâche de trouver la bonne erreur E_j s'étant produite à partir des syndromes mesurés est celle du décodeur. Dans ce chapitre, nous étudions la performance d'un code de surface de distance $d = 3$ et nous nous contentons d'utiliser un algorithme standard pour réaliser le décodage des syndromes [57].

Le code de surface [58] est utilisé et étudié par de nombreux groupes de recherche en raison de l'excellent compromis qu'il offre entre ses contraintes expérimentales et sa performance logique [16]. En effet, la mesure des stabilisateurs du code de surface se décompose en un circuit n'utilisant que 4 opérations à deux qubits par stabilisateur (des CNOT ou CZ) et ces qubits peuvent être connectés de façon locale à seulement 4 voisins sur une grille cartésienne en deux dimensions. Même sous ces contraintes de connectivité locale, le code de surface possède une très bonne protection contre les erreurs survenant dans le circuit quantique, comme nous le verrons dans l'article de la section 3.4. Cette performance peut également être quantifiée à l'aide de simulations numériques du code de surface en présence d'erreurs de dépolarisation associées à chacune des opérations réalisées. Avec ce modèle où chaque opération produit une erreur de Pauli avec une probabilité p , il existe un seuil d'erreur p_{seuil} sous lequel la probabilité d'obtenir une erreur logique diminue de façon exponentielle avec la taille du code de surface telle que

$$p_L \propto \left(\frac{p}{p_{\text{seuil}}} \right)^{\frac{d+1}{2}}, \quad (3.3)$$

où la constante de proportionnalité est environ 0.03 et le taux d'erreur seuil est environ de 1 % [16]. Dans l'objectif d'atteindre $p_L \approx 10^{-10}$, il est alors nécessaire d'opérer sous le seuil d'erreur du code et d'utiliser un code avec une grande distance, ce qui nécessite $n = 2d^2 - 1$ qubits physiques. Étant donné la suppression exponentielle des erreurs fournie par le code, on comprend que toute amélioration physique des opérations réalisées (toute réduction de la probabilité p) peut avoir un grand impact sur l'empreinte physique qu'il sera nécessaire d'utiliser afin de concevoir un ordinateur quantique universel qui est utile.

3.2 Modèles de bruit en QEC

Le développement des codes stabilisateurs, tout comme la présence du seuil d'erreur obtenu numériquement, est basé sur le fait que les qubits sont réellement des systèmes à deux niveaux et que les erreurs provenant des différentes opérations sont purement stochastiques et indépendantes. Ces hypothèses ne sont pas respectées en pratique en utilisant des transmons comme qubits. En effet, de nombreuses imperfections peuvent causer des déviations importantes au modèle de bruit stochastique local, telles que la fuite de niveaux, les erreurs cohérentes produisant des portes unitaires imparfaites, le couplage résiduel continuellement actif entre les transmons ainsi que l'ensemble des opérations logiques ayant un effet non trivial sur les qubits avoisinants, par exemple suite à une diaphonie des signaux de contrôle. En plus de l'ingénierie précise de toutes les composantes et opérations du système, la simulation numérique de l'impact de ces différentes imperfections devient cruciale à notre compréhension et à notre capacité à opérer un code de correction d'erreur à grande échelle. Je survole ici les différentes approches utilisées afin de simuler la réalisation d'un code de surface.

Tel qu'illustré schématiquement à la Fig. 3.1, il existe trois approches principales afin de simuler le bruit provenant d'un circuit de mesure des stabilisateurs. Tout d'abord, l'approche la plus simple et la plus efficace numériquement consiste à utiliser un canal de dépolarisation afin d'appliquer des erreurs de Pauli de façon stochastique après ou avant chacune des opérations. Le bruit prend alors la forme

$$\mathcal{E}_{\text{depol.}}(\hat{\rho}_{1q}) = (1 - p_{\Sigma})\mathbb{1}\hat{\rho}_{1q}\mathbb{1} + p_x \hat{\sigma}_x \hat{\rho}_{1q} \hat{\sigma}_x + p_y \hat{\sigma}_y \hat{\rho}_{1q} \hat{\sigma}_y + p_z \hat{\sigma}_z \hat{\rho}_{1q} \hat{\sigma}_z, \quad (3.4)$$

où $p_{\Sigma} = p_x + p_y + p_z$. Le qubit subit donc une erreur de bit avec une probabilité p_x , une erreur de phase avec une probabilité p_z ou les deux avec une probabilité p_y . Ce canal est également généralisé à deux qubits suivant les opérations CNOT ou CZ et dans ce cas les 15 probabilités associées aux 15 générateurs non triviaux du groupe $\mathcal{P}_2$ sont typiquement choisies comme étant les mêmes, souvent par manque d'information concernant les erreurs unitaires réellement présentes. Il s'agit cependant d'un choix significatif étant donné que certains générateurs (9/15) correspondent à des erreurs de poids deux (par exemple $\hat{\sigma}_x \otimes \hat{\sigma}_y$), lesquels sont beaucoup plus dommageables pour le code que les erreurs de poids un (6/15, par exemple $\mathbb{1} \otimes \hat{\sigma}_z$). Au chapitre 5, nous développerons d'ailleurs un protocole de caractérisation précis de ces erreurs cohérentes, lequel pourrait être utilisé directement afin d'obtenir des probabilités d'erreurs de Pauli plus représentatives de l'expérience.

Étant donné que les circuits de codes stabilisateurs utilisent des portes de Clifford et

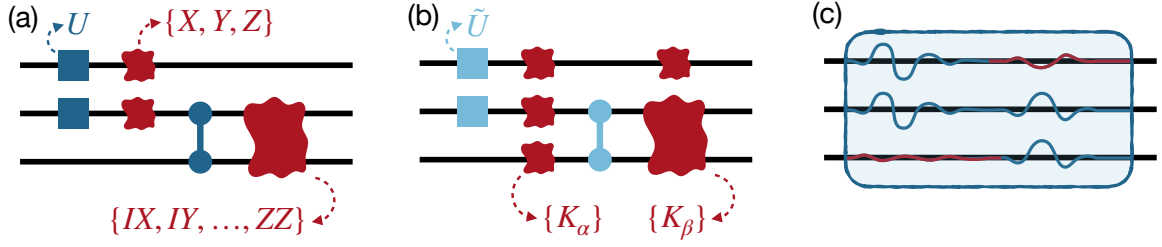


FIGURE 3.1 – Différentes approches pour simuler le bruit associé à la réalisation d’un circuit quantique. (a) Le modèle le plus répandu en correction d’erreur quantique ne considère que des erreurs de Pauli, formant ainsi un canal dépolarisant (rouge) qui suit l’application des portes quantiques parfaites (bleu). Lorsque les opérations unitaires sont des portes de Clifford, un tel circuit comportant des milliers de qubits peut être simulé de façon efficace sur un ordinateur classique [60]. (b) Les canaux de bruits affectant quelques qubits peuvent être généralisés à des opérateurs de Kraus arbitraires. Des algorithmes classiques sont disponibles afin de simuler un tel modèle général impliquant jusqu’à environ 40 qubits [61]. (c) Au-delà d’opérations discrètes n’affectant que quelques qubits de façon locale, il est possible de simuler l’ensemble des dynamiques temporelles d’un circuit quantique à partir des signaux de contrôle réellement utilisés dans l’expérience. Cette méthode utilisant par exemple une équation maîtresse peut difficilement être appliquée à plus d’une vingtaine de qubits.

que le modèle de bruit n’applique que des opérations de Pauli, en vertu du théorème de Gottesman-Knill, il est possible d’utiliser une représentation classique très efficace afin de simuler de manière exacte ces circuits en un temps polynomial sur un ordinateur classique [59]. Cette approche est d’ailleurs utilisée dans la librairie en code source ouvert `Stim`, laquelle permet de simuler efficacement des codes de surface comportant des milliers de qubits correspondant à des distances $d > 100$ [60].

La seconde approche consiste à généraliser le bruit ainsi que les opérations unitaires à des canaux quantiques arbitraires représentés à l’aide d’un ensemble d’opérateurs de Kraus, une approche aussi appelée la représentation de la somme d’opérateurs [62]. Dans ce cas, le canal quantique de chaque porte logique est représenté par

$$\mathcal{E}_{\text{Kraus}}(\hat{\rho}_l) = \sum_{\alpha} \hat{K}_{\alpha} \hat{\rho}_l \hat{K}_{\alpha}^{\dagger}, \quad (3.5)$$

où l indique le nombre de qubits impliqués et $\sum_{\alpha} \hat{K}_{\alpha}^{\dagger} \hat{K}_{\alpha} = 1$ pour toute opération préservant la trace de la matrice densité. Dans cette représentation, une transformation unitaire ne possède qu’un seul opérateur de Kraus $\hat{K} = \hat{U}$, de sorte que $\mathcal{U}(\hat{\rho}) = \hat{U} \hat{\rho} \hat{U}^{\dagger}$. En pratique, les portes logiques sont souvent simulées en utilisant la décomposition approximative $\mathcal{E} = \mathcal{A} \circ \mathcal{U}$, où la transformation unitaire (potentiellement différente de l’opération cible) $\mathcal{U}$ est suivie par un canal de relaxation et de déphasage du qubit $\mathcal{A}$ [63]. Cette approche permet de grandement limiter le nombre d’opérateurs de Kraus nécessaires.

La décomposition par opérateurs de Kraus permet donc de discrétiser le circuit quantique

en une séquence d'opérations dont l'effet est obtenu à l'aide de plusieurs multiplications de matrices avec l'équation Éq. (3.5). Étant donné que le nombre d'opérateurs nécessaires pour décrire un canal composé de l qubits peut aller jusqu'à 4^l , des canaux impliquant seulement $l = 1$ ou $l = 2$ sont typiquement utilisés afin de restreindre les ressources numériques nécessaires [63]. De ce fait, cette approche ne prend typiquement pas en compte les interactions parasites entre les qubits, lesquelles nécessiteraient d'utiliser un support sur de nombreux qubits (par exemple $l = 8$ considérant les qubits voisins lors des opérations CZ [64]). Avec de telles approximations, il est possible de simuler jusqu'à environ 40 qubits à l'aide de la librairie en code source ouvert qsim [61].

Enfin, la troisième et dernière approche que je considère pour modéliser l'expérience de correction d'erreur est l'utilisation d'une équation maîtresse (EM). Tel que présenté au chapitre dernier, cette équation différentielle prend la forme

$$\mathcal{E}_{\text{EM}}(\hat{\rho}_l) = \int dt \mathcal{L} \hat{\rho}_l, \quad (3.6)$$

où $l = 17$ pour un code de surface de distance trois. Ce modèle nous permet de simuler l'évolution temporelle du système de manière continue et d'inclure des interactions impliquant de nombreux qubits, comme nous le verrons dans la section qui suit avec la simulation des interactions parasites ZZ entre tous les qubits voisins. Nous utilisons une telle équation maîtresse pour simuler l'expérience étant donné que ce problème de diaphonie quantique est présent dans tout dispositif supraconducteur et que les erreurs corrélées impliquant de nombreux qubits peuvent briser les hypothèses d'une correction d'erreur quantique réussie avec un code de distance trois.

3.3 Méthodologie

Dans cette section, je décris la méthodologie utilisée pour simuler l'expérience de correction d'erreur quantique réalisée à l'ETH Zurich [1], laquelle implique 17 qubits. Ces simulations sont basées sur les caractéristiques mesurées directement sur le dispositif, telles que les temps de cohérence de chacun des qubits à différents points d'opération, les erreurs de mesure et le couplage entre les qubits voisins. Étant donné que nous ne sommes pas en mesure d'inclure l'ensemble des sources de bruit du dispositif, les résultats de ces simulations fournissent une limite supérieure à la performance réalisable avec ce dispositif. Par exemple, nous n'incluons pas les erreurs cohérentes de contrôle des portes logiques, les fuites de niveaux des transmons, les pertes de population associées au couplage avec des

défauts à deux niveaux dans les matériaux ni le déphasage des qubits voisinant produit par la mesure dispersive du transmon.

Dans les sous-sections qui suivent, je décris en détail notre approche de simulation ainsi que les différentes approximations que nous avons réalisées afin de rendre la description juste d'une telle expérience possible en pratique. Je prends soin de justifier chacune de ces approximations en plus de démontrer numériquement la validité de plusieurs d'entre elles.

3.3.1 Hamiltonien du système

Le système est composé de 17 transmons de fréquence variable couplés capacitivement à 4 voisins dans une structure de grille cartésienne. L'Hamiltonien décrivant le dispositif prend alors la forme

$$\hat{H}_{\text{sys}} = \sum_{k=1}^{17} 4E_{C,k} \hat{n}_k^2 - E_{J,k}(\Phi_k) \cos(\hat{\phi}_k - \varphi_{0,k}) + \sum_{k < k'} \hbar g_{k,k'} \hat{n}_k \hat{n}_{k'}, \quad (3.7)$$

où $\hbar g_{k,k'}$ est l'amplitude du couplage capacitif, Φ_k est le flux externe appliqué dans la boucle supraconductrice du transmon k et $\tan(\varphi_{0,k}) = d \tan(\pi \Phi_k / \Phi_0)$ avec Φ_0 le quantum de flux et d l'asymétrie entre les deux jonctions Josephson utilisées [27]. Cet Hamiltonien dépend également du temps lors de la réalisation de portes logiques, par exemple avec des flux externes variant sur chacun des qubits $\Phi_k(t)$ lors des portes CZ et avec un terme de conduite de la forme $A_k(t) \cos(\omega_{d,k}t + \phi_{d,k}) \hat{n}_k$ pour les portes à un qubit.

La simulation des dynamiques associées à un tel Hamiltonien est très ardue, même en ne considérant aucun bruit et en utilisant l'Éq. (2.26). En effet, l'utilisation approximative de m niveaux par transmon nécessite un espace d'Hilbert de taille m^{17} , ce qui représente 130 millions d'états pour $m = 3$. Le simple stockage de cet Hamiltonien nécessite plus de 130'000 TB de mémoire en format dense, un requis simplement inaccessible en pratique¹. Il devient alors nécessaire d'utiliser des approximations afin de simplifier le problème numérique.

L'approximation la plus importante que nous utilisons consiste à ne considérer que les deux premiers niveaux du transmon ($m = 2$). Cette approximation est partiellement justifiée par le fait que l'expérience rejette, avec un processus de sélection a posteriori, toutes

1. Étant donné que la matrice hamiltonienne n'est pas complètement dense, il est possible d'utiliser une représentation creuse et de diminuer ce coût en mémoire par un facteur d'amplitude considérable (typiquement plus de 100, une valeur qui dépend du ratio du nombre d'entrées non nulles dans la matrice). Le requis en mémoire présenté est le plus pessimiste et sert principalement à illustrer la complexité du problème que l'on souhaite résoudre.

les réalisations où une fuite de niveau est détectée. En effet, chaque mesure dispersive réalisée discrimine également l'état $|2\rangle$ des transmons ce qui permet de détecter directement la présence d'une ou de plusieurs fuites de niveau et donc de limiter la grande majorité des expériences conservées à des dynamiques restreintes au sous-espace de calcul. Sous ce modèle d'un ensemble de modes à deux niveaux, le modèle effectif du système devient

$$\hat{H}_{\text{eff}}(t)/\hbar = \sum_{k=1}^{17} \frac{\omega_k[\Phi_k(t)]}{2} \hat{\sigma}_k^z + \sum_{k < k'} \xi_{k,k'}[\Phi_k(t), \Phi_{k'}(t)] \hat{\sigma}_k^z \hat{\sigma}_{k'}^z + \sum_k \hat{H}_{d,k}(t), \quad (3.8)$$

où $\omega_k[\Phi_k(t)]$ est la fréquence du qubit k , $\hat{H}_{d,k}(t)$ sont des Hamiltoniens associés au contrôle des portes logiques que nous verrons plus bas et $\xi_{k,k'}[\Phi_k(t), \Phi_{k'}(t)]$ est la force de l'interaction ZZ (ou l'interaction croisée de Kerr) provenant du couplage capacitif entre les transmons. Ce dernier terme dépend fortement de la différence de fréquence entre ces qubits, d'où la dépendance aux flux externes. Nous incluons cette dépendance puisque la fréquence des qubits varie fortement (de l'ordre de 1 GHz) afin d'activer les opérations à deux qubits, causant des changements importants dans la force de l'interaction ZZ parasite entre les qubits impliqués dans les portes et leurs voisins. Grâce à la conception soignée de la géométrie et des nombreux paramètres du dispositif comme la forme des signaux et les fréquences d'opération utilisées, il s'avère que $\hat{H}_{\text{eff}}$ génère une bonne description des dynamiques cohérentes du dispositif et que les nombreuses interactions négligées avec cet Hamiltonien sont largement minimisées. Nous en verrons la démonstration à la section 3.4.

Afin de décrire l'effet d'un environnement de température effective nulle, nous incluons les processus de relaxation et de déphasage des transmons de manière phénoménologique à l'aide de l'équation maîtresse

$$\frac{d}{dt} \hat{\rho} = -\frac{i}{\hbar} [\hat{H}_{\text{eff}}(t), \hat{\rho}] + \sum_k \gamma_{1,k} \mathcal{D}[\hat{\sigma}_k^-] \hat{\rho}(t) + \sum_k \gamma_{\varphi,k}[\Phi_k(t)] \mathcal{D}[\hat{\sigma}_k^z] \hat{\rho}(t), \quad (3.9)$$

où $\gamma_{1,k}$ est le taux de relaxation du qubit k et $\gamma_{\varphi,k}[\Phi_k(t)]$ son taux de déphasage, lequel change significativement lors de l'exécution de portes CZ. Tous les paramètres entrant dans l'Hamiltonien $\hat{H}_{\text{eff}}(t)$ et dans cette équation maîtresse ont été mesurés expérimentalement pour chaque qubit et sont utilisés directement dans nos simulations, incluant le taux de déphasage des qubits pendant les portes CZ et pendant les autres opérations.

3.3.2 Approche numérique

Afin de simuler l'expérience de correction d'erreur quantique, il est nécessaire d'intégrer l'équation différentielle de l'Éq. (3.9) pour la séquence de portes réalisant jusqu'à 16 rondes

de la mesure des stabilisateurs du code de distance 3, chaque ronde durant $1.1 \mu\text{s}$. Afin de rendre cette simulation possible, nous utilisons la méthode des fonctions d'onde de Monte-Carlo [65, 66, 67], disponible avec la librairie QuTiP [68]. Le principal avantage de cette approche est qu'elle n'utilise que des états purs (fonctions d'ondes), lesquels sont des vecteurs de taille m^{17} plutôt que les matrices densité de taille $m^{17} \times m^{17}$ utilisées dans l'équation maîtresse. La réduction majeure des besoins en mémoire rend cette approche significativement plus facile à utiliser pour nos besoins. Tel qu'expliqué plus bas, cette approche stochastique nécessite cependant un coût d'échantillonnage additionnel.

La méthode de Monte-Carlo consiste à faire évoluer un état pur $|\psi(t)\rangle$ avec un Hamiltonien non hermitien de la forme

$$\hat{H}_{\text{nh}} = \hat{H}_{\text{eff}}(t) - \frac{i\hbar}{2} \sum_l \hat{c}_l^\dagger \hat{c}_l, \quad (3.10)$$

où $\hat{c}_l$ sont les opérateurs de dissipation de l'équation maîtresse originale. L'évolution temporelle sous cet Hamiltonien pour une durée infinitésimale dt mène à une réduction de la norme carrée de l'état quantique initialement pur de

$$dp = \sum_l dt \langle \psi(t) | \hat{c}_l^\dagger \hat{c}_l | \psi(t) \rangle \ll 1. \quad (3.11)$$

Un processus stochastique est alors utilisé afin d'obtenir l'état au nouveau temps $t + dt$. Avec une probabilité de $1 - dp$, la fonction d'onde $|\psi(t + dt)\rangle$ est renormalisée (par le facteur $1 - dp$) afin de retrouver un état pur.² Avec une probabilité restante de dp , la fonction d'onde est plutôt obtenue en appliquant un *saut quantique* à l'état $|\psi(t)\rangle$ tel que

$$|\psi(t + dt)\rangle = \frac{\hat{c}_l |\psi(t)\rangle}{\langle \psi(t) | \hat{c}_l^\dagger \hat{c}_l | \psi(t) \rangle^{1/2}}, \quad (3.12)$$

où le choix de l'opérateur de saut $\hat{c}_l$ est choisi de manière stochastique à partir des probabilités

$$p_l = \frac{\langle \psi(t) | \hat{c}_l^\dagger \hat{c}_l | \psi(t) \rangle}{\sum_{l'} \langle \psi(t) | \hat{c}_{l'}^\dagger \hat{c}_{l'} | \psi(t) \rangle}, \quad (3.13)$$

de chaque opérateur de saut.

Le résultat d'une telle évolution stochastique est appelé une trajectoire quantique et correspond à l'évolution d'un état pur interrompu par des sauts discrets dont la fréquence et le caractère sont dictés par la forme de l'équation maîtresse originale. Un avantage de

2. Afin d'éviter plusieurs opérations numériques, je note qu'en pratique la fonction d'onde n'est pas renormalisée à chaque pas temporel, mais seulement après chaque saut quantique et à la toute fin de l'évolution.

cette approche est que chacune des trajectoires quantiques peut être calculée de façon indépendante, de sorte que de nombreuses solutions peuvent être obtenues en parallèle en utilisant les différents fils d'exécution de l'ordinateur classique.

Avec un nombre suffisant de trajectoires, il est possible de simplement moyenner les résultats obtenus afin d'obtenir la solution à l'équation maîtresse originale $\mathbb{E}[|\psi(t)\rangle\langle\psi(t)|] = \hat{\rho}(t)$. La valeur attendue des observables d'intérêt peut être obtenue de façon similaire avec

$$\mathbb{E}[\langle\psi(t)|\hat{O}|\psi(t)\rangle] = \text{Tr}[\hat{\rho}(t)\hat{O}]. \quad (3.14)$$

Dans l'objectif de nous assurer d'utiliser suffisamment de trajectoires quantiques pour que l'Éq. (3.14) soit satisfaite avec le niveau de précision requis, nous analysons la convergence des valeurs attendues de tous les stabilisateurs du code et des opérateurs logiques du code en fonction du nombre de trajectoires utilisées. Nous trouvons que cinquante mille trajectoires suffisent afin d'estimer ces valeurs attendues avec moins de 1 % d'incertitude relative. Je me suis également assuré que les résultats obtenus convergeaient vers les solutions exactes. Pour ce faire, j'ai comparé les matrices densité obtenues en solutionnant avec l'approche Monte-Carlo à celles obtenues en solutionnant directement l'équation maîtresse pour de plus petits systèmes comportant jusqu'à 7 qubits (ce qui correspond à un code de surface de distance 2 [69]). En utilisant des tolérances d'erreur qui sont de l'ordre de la précision numérique lors de la résolution de l'équation maîtresse, j'ai pu ajuster les tolérances et paramètres d'échantillonnage de l'algorithme Monte-Carlo afin d'assurer une convergence vers la bonne solution sur le plus grand système possédant 17 qubits.

3.3.3 Opérations logiques

Je décris maintenant les 3 opérations logiques réalisées dans l'expérience de correction d'erreur quantique : les portes à un qubit, à deux qubits et la mesure projective des qubits. Tel qu'illustré à la Fig. 3.2, ces trois opérations peuvent ensuite être combinées afin de former un cycle complet de correction d'erreurs quantiques à l'aide du code de surface.

Les portes à un qubit de l'expérience sont réalisées en $T_{SQ} = 40$ ns et sont représentées par les unitaires $R_y(-\pi/2)$, $R_y(\pi/2)$ et $R_y(\pi)$. Cette durée relativement longue est utilisée afin de minimiser la diaphonie classique des opérations alors que les contrôles appliqués sur un qubit peuvent influencer l'état des autres qubits (je discuterai de cet effet plus en détail à la section 5.5). De ce fait, les erreurs cohérentes des portes à un qubit sont minimales et la fidélité de ces opérations est principalement limitée par la cohérence des transmons. Cette

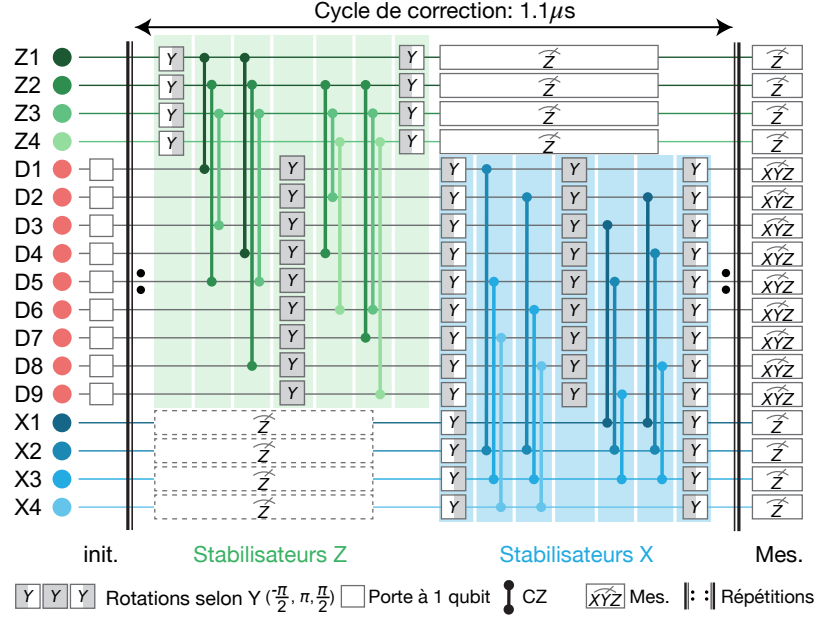


FIGURE 3.2 – Circuit quantique réalisant la stabilisation d’un code de surface de distance trois sur 17 qubits. Les qubits en rouge sont les qubits du code (D1 à D9) et les quatre qubits auxiliaires en vert (bleu) sont utilisés pour mesurer les stabilisateurs de type Z (X).

caractéristique expérimentale nous permet d’éviter de simuler l’évolution précise des qubits lors des portes à un qubit où des impulsions micro-ondes dépendantes du temps $\Omega_k(t)$ sont appliquées. En effet, on se contente d’utiliser un Hamiltonien constant de la forme

$$\hat{H}_{d,k}^{SQ} = \frac{\theta}{2T_{SQ}} \hat{\sigma}_k^y \quad (3.15)$$

pendant $T_{SQ} = 40$ ns afin de réaliser l’unitaire cible sur les qubits, où $\theta \in \{-\pi/2, \pi/2, \pi\}$. Ainsi, le bruit associé à ces opérations provient seulement de l’interaction ZZ parasite entre les qubits et de la relaxation et du déphasage des transmons. Ce modèle produit une fidélité de porte moyenne pour l’ensemble des opérations et l’ensemble des qubits de 0.08 %, ce qui correspond de près à la valeur de 0.09 % mesurée expérimentalement. J’ai également effectué des simulations de la détection d’erreur quantique avec 7 des qubits du dispositif utilisé ici (un code de surface de distance 2 [69]) avec et sans l’utilisation d’une telle approximation de contrôle constant et je n’ai obtenu aucune différence significative sur la performance du code.

Les portes à deux qubits sont des portes CZ et sont réalisées dans l’expérience en appliquant un flux magnétique à travers la boucle supraconductrice de chacun des transmons couplés capacitivement de sorte que les niveaux $|11\rangle$ et $|20\rangle$ soient en résonance pour une

durée précise [70, 71]. Étant donné notre traitement des transmons en tant que systèmes à deux niveaux, il ne nous est pas possible de simuler directement les dynamiques temporelles associées à ces opérations. Encore une fois, la sélection a posteriori des données de l'expérience ne contenant aucune fuite de niveaux nous permet de nous concentrer uniquement sur les dynamiques dans le sous-espace de calcul. De façon similaire aux portes à un qubit, nous utilisons alors un Hamiltonien constant

$$\begin{aligned}\hat{H}_{d,k,k'}^{\text{CZ}} &= \left(\frac{\pi}{T_{\text{CZ}}} - \xi_{k,k'} \right) \frac{\mathbb{1}_k - \hat{\sigma}_k^z}{2} \otimes \frac{\mathbb{1}_{k'} - \hat{\sigma}_{k'}^z}{2} \\ &= \left(\frac{\pi}{T_{\text{CZ}}} - \xi_{k,k'} \right) |1\rangle\langle 1|_k \otimes |1\rangle\langle 1|_{k'},\end{aligned}\tag{3.16}$$

pour les durées utilisées d'environ $T_{\text{CZ}} = 100$ ns, alors qu'on vérifie facilement que $\text{CZ} = e^{-i\pi|11\rangle\langle 11|}$.

En plus de cette modification à l'Hamiltonien pour générer l'unitaire du CZ, il est nécessaire de prendre en compte que la fréquence des qubits change significativement lors de ces opérations. En effet, les transmons sont généralement opérés au point où leur sensibilité aux variations de flux magnétique externe est minimale (ce qui correspond à environ 6 GHz pour les qubits ancillaires et 4 GHz pour les qubits du code), mais doivent être déplacés à des flux où leur taux de déphasage est plus élevé afin d'activer l'interaction du CZ (à des fréquences d'environ 5 GHz). À ces fréquences, l'interaction ZZ parasite avec les qubits voisins est également accrue en raison d'un décalage en fréquence beaucoup plus faible entre les qubits qui ne devraient pas interagir.

Nous prenons en compte ces deux effets dans nos simulations en utilisant directement les fréquences d'interactions utilisées par les 24 portes CZ du code afin de calculer les interactions ZZ modifiées avec tous les qubits voisins de ces opérations, en plus d'utiliser les taux de cohérence (T_2^*) mesurés expérimentalement à la fréquence d'interaction de chaque transmon et de chaque porte. Ces temps de cohérence sont en moyenne trois fois plus courts que lorsqu'aucune porte CZ n'est réalisée sur le transmon. Avec ce modèle, nous obtenons une erreur de porte moyenne de 0.9 % en simulation alors que l'erreur moyenne évaluée expérimentalement est de 1.5 %. Je présenterai à la section 3.3.5 une modification au modèle de bruit permettant d'obtenir la même infidélité moyenne que dans l'expérience, ainsi que l'effet de cette modification sur la performance simulée du code de surface.

Enfin, les mesures dispersives des transmons sont simulées de manière phénoménologique en faisant évoluer les qubits à l'aide de notre équation maîtresse pour la durée totale de la mesure ($T_{\text{mesure}} \approx 400$ ns), avant d'effectuer une projection idéale de l'état du système

en suivant la règle de Born

$$|\psi\rangle \rightarrow \frac{\Pi_k^j |\psi\rangle}{\sqrt{p_k^j}}, \quad (3.17)$$

où Π_k^j est le projecteur dans la base de calcul du qubit k , $j \in \{0, 1\}$ correspond à l'état mesuré du qubit et $p_k^j = \langle \psi | \Pi_k^j | \psi \rangle$ est la probabilité de mesurer le qubit k dans l'état j . Par la suite, nous incluons les erreurs de mesures provenant d'une discrimination imparfaite des états en inversant le bit obtenu j en utilisant la matrice d'affection de mesure qui a été caractérisée expérimentalement pour chaque qubit, c'est-à-dire en utilisant les probabilités d'erreur $P(0|1)$ et $P(1|0)$.

Rassemblant nos modèles des portes logiques et des mesures projectives, il nous est possible de simuler exactement la séquence d'opérations utilisée dans l'expérience de correction d'erreur quantique, laquelle inclut d'ailleurs plusieurs durées tampons afin de synchroniser les opérations possédant des durées légèrement différentes. Utilisant l'algorithme de Monte-Carlo décrit plus haut, nous sommes alors en mesure de simuler l'évolution du système de 17 qubits pour 16 cycles de correction d'erreur, ce qui correspond environ à 20 μ s. Cependant, les ressources numériques en mémoire et en temps de calcul pour obtenir suffisamment de trajectoires quantiques sont élevées et nous souhaitons simuler cette expérience avec un ensemble de paramètres différents afin d'étudier la performance du code lors de diverses améliorations au dispositif. Nous utilisons alors une simplification provenant de la structure du circuit quantique réalisé afin de grandement accélérer ces simulations.

3.3.4 Modèle effectif 9+4

Dans l'expérience, les stabilisateurs de type Z et X sont mesurés séquentiellement afin de réduire le nombre maximal de portes CZ réalisées en parallèle à seulement trois. Ce choix est motivé par le fait que le couplage capacitif fixe entre les transmons crée des interactions parasites significatives lorsque leur décalage en fréquence est réduit [64]. Nous sommes en mesure d'exploiter cette structure afin de réduire le nombre de qubits dans l'espace d'Hilbert simulé de 17 à 13 tout en obtenant essentiellement les mêmes dynamiques. Nous surnomons cette modélisation effective 9+4 en référence à l'espace d'Hilbert utilisé avec les 9 qubits du code et un registre de seulement 4 qubits auxiliaires. Voyons plus en détail comment ce modèle est construit.

Tel qu'illustré schématiquement à la Fig. 3.3(a), le circuit quantique permettant de réaliser le code de surface de distance trois peut être séparé en deux moitiés où toutes les opérations à un et deux qubits n'agissent que sur les 9 qubits du code (rouge) et 4 qubits

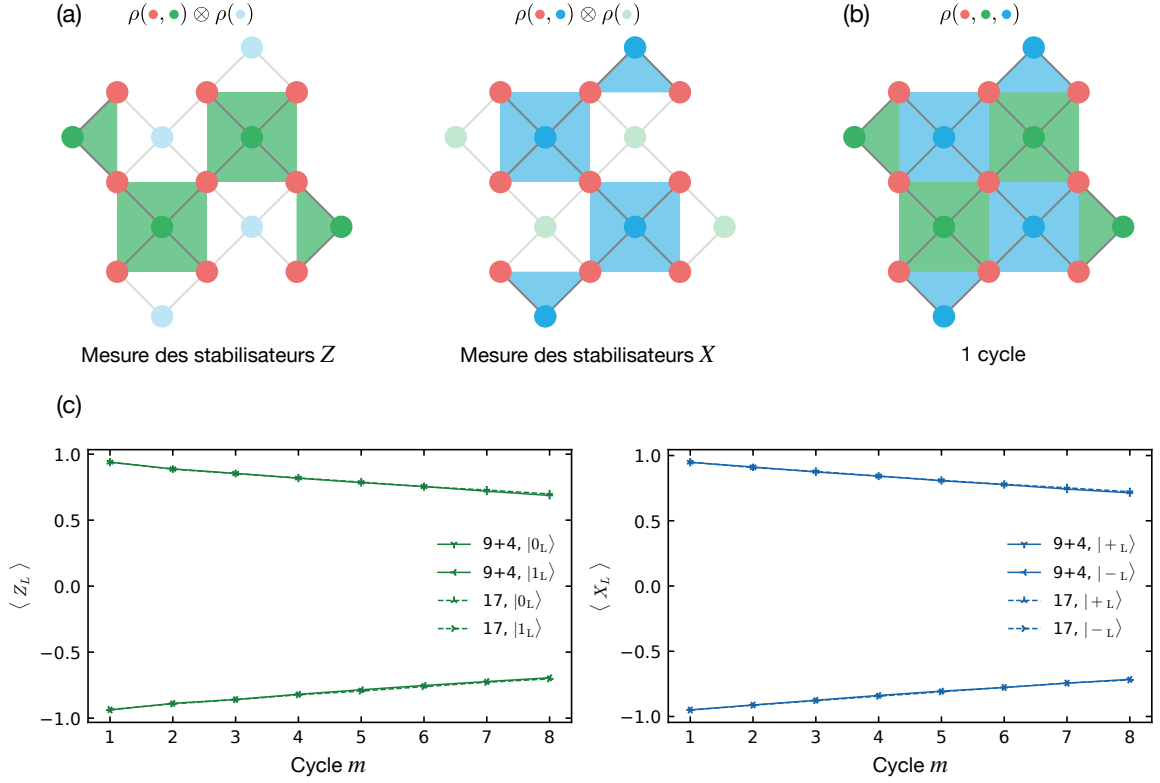


FIGURE 3.3 – Illustration du modèle effectif 9+4. (a) La mesure des stabilisateurs Z (X) est simulée en retirant de l’espace d’Hilbert utilisé, à l’aide d’une trace partielle, les qubits ancillaires de type X (Z). La matrice densité $\rho(\bullet)$ inclut les qubits associés au code de couleur de la figure. Seulement 13 qubits (9 qubits du code et 4 qubits ancillaires) sont nécessaires pour décrire les opérations appliquées sur le système à tout moment. (b) Le modèle complet inclut l’ensemble des 17 qubits, tout au long de la simulation. (c) Résultats de simulation de l’expérience de stabilisation du code de surface impliquant les 17 qubits, en utilisant le modèle effectif (9+4) et le modèle total (17). On observe un excellent accord des temps de vie logique T_1 et T_2 entre les deux modèles de simulations.

ancillaires de type Z (vert) ou X (bleu). Pendant ce temps, la mesure dispersive des qubits ancillaires de l’autre type, laquelle dure 400 ns des 550 ns du demi-cycle, projette ces qubits dans un état pur séparable. Par exemple pour les quatre qubits ancillaires de type X, dénotés par k'' , cet état prend la forme

$$|X_n\rangle = \otimes_{k''} |j_{k''}\rangle, \quad (3.18)$$

où $j_{k''} \in \{0, 1\}$. En supposant que les qubits ancillaires de type X restent approximativement dans un état séparable pendant toute la durée du demi-cycle, il est possible d’éliminer ces modes de l’Hamiltonien lors de l’évolution temporelle du système, mais de tout de même considérer leur effet sur les 13 qubits restants. Pour ce faire, nous incluons de manière effective l’interaction ZZ provenant de l’état dans lequel se trouvent les qubits éliminés. Nous tenons également compte du temps de vie fini des qubits ancillaires $T_{1,k''}$ en modifiant

l'Hamiltonien effectif de l'Éq. (3.8) de sorte à obtenir

$$\begin{aligned} \hat{H}_{9+4}(t)/\hbar = & \sum_{k=1}^{13} \frac{1}{2} \left(\omega_k[\Phi_k(t)] + \sum_{k''=1}^4 \xi_{k,k''}[\Phi_k(t)] \frac{1 - \langle \hat{\sigma}_{k''}^z \rangle}{2} e^{-t/T_{1,k''}} \right) \hat{\sigma}_k^z \\ & + \sum_{k < k'} \xi_{k,k'}[\Phi_k(t), \Phi_{k'}(t)] \hat{\sigma}_k^z \hat{\sigma}_{k'}^z + \sum_k \hat{H}_{d,k}(t), \end{aligned} \quad (3.19)$$

où $\langle \hat{\sigma}_{k''}^z \rangle \in \{-1, 1\}$ est la valeur attendue du qubit ancillaire éliminé directement au début du demi-cycle, sachant que l'interaction ZZ avec un qubit éliminé par une trace partielle crée simplement un décalage de la fréquence du qubit restant. Nous prenons en compte la dépendance temporelle de ce décalage lors des portes CZ (avec la dépendance temporelle $\xi_{k,k''}[\Phi_k(t)]$) et à mesure que les qubits ancillaires relaxent (avec le terme $e^{-t/T_{1,k''}}$).

À la suite de l'évolution temporelle pour $T = 550$ ns sous $\hat{H}_{9+4}(t)$, les qubits ancillaires de type Z sont mesurés de manière projective et l'état simulé est remplacé en utilisant

$$|\psi(T)\rangle_9 |Z_n\rangle_4 \rightarrow |\psi(T)\rangle_9 |X_n\rangle_4, \quad (3.20)$$

où $|\psi(T)\rangle_9$ représente l'état des 9 qubits du code après avoir projeté l'état quantique des 13 qubits. Plutôt que d'utiliser directement l'état quantique des qubits de type X mesuré au tout début du demi-cycle, nous tenons compte de leur durée de vie en changeant les états $|1\rangle$ vers $|0\rangle$ avec une probabilité $1 - \exp(-T_{\text{demi-cycle}}/T_{1,k''})$. La simulation du second demi-cycle, mesurant les stabilisateurs de type X, est effectuée de la même manière que le premier.

La Fig. 3.3(c) illustre la validité du modèle 9+4 en comparant directement la performance du code de surface obtenue avec ce modèle à celle obtenue avec le modèle simulant les 17 qubits. Tel que présenté dans l'article suivant à la section 3.4, la performance du code est évaluée en préparant un état propre du code et en le stabilisant pour plusieurs cycles de correction d'erreur, ici jusqu'à 8 cycles. Un ajustement à la décroissance exponentielle observée permet d'extraire les temps de vie logique $T_{1,L}$ et $T_{2,L}$ du qubit encodé. Nos résultats présentent un excellent accord entre les deux modèles, avec un désaccord qui est contenu dans l'erreur associée à l'ajustement exponentiel. Bien qu'il ne permette que de simuler un nombre limité de processus de bruit, je note que notre modèle effectif va au-delà des approches précédentes [72, 73] dans le traitement des erreurs corrélées (interactions ZZ) et de la dissipation agissant de manière continue pendant les différentes opérations.

Le modèle 9+4 permet de considérablement réduire les besoins numériques de simulation par rapport à l'utilisation de 17 qubits. Avec notre approche Monte-Carlo et ordinateur performant (vitesse d'horloge de 3 GHz), le modèle de 17 qubits requiert environ 25 s par

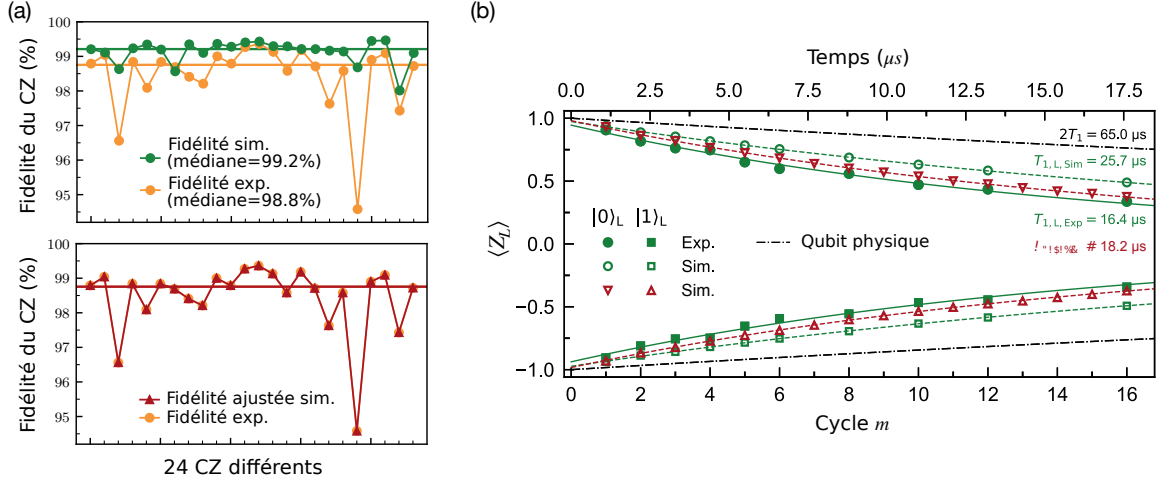


FIGURE 3.4 – Ajustement du modèle de bruit. (a) Les portes à deux qubits CZ simulées à partir de notre modèle effectif et des paramètres de cohérence de tous les transmons du dispositif (vert) produisent des fidélités qui sont plus élevées que celles observées expérimentalement (orange). Il est possible d’ajuster le temps de déphasage des qubits (T_1^*) lors de ces opérations afin de reproduire les fidélités mesurées avec précision (rouge). (b) Reproduction de la figure 4a de la Réf. [1] où nous ajoutons en rouge les résultats provenant du modèle où les fidélités des portes à deux qubits ont été ajustées.

trajectoire quantique et par cycle de correction d’erreur, alors que le modèle 9+4 nécessite environ 2 s. De même, les besoins en mémoire sont réduits de 6 GB par trajectoire quantique à seulement 1 GB. Afin d’accumuler suffisamment de statistiques avec un ordinateur possédant 64 fils d’exécution parallèle, notre modèle effectif requiert environ 7 h pour simuler la plus longue expérience réalisée alors que l’approche à 17 qubits nécessite 87 h. Étant donné ces ressources d’envergure et le fait que nous devons simuler de nombreuses réalisations avec différents paramètres expérimentaux, on comprend que toute simplification justifiée est la bienvenue pour simuler de grands systèmes quantiques.

3.3.5 Modèle de bruit ajusté

La plus grande divergence entre notre modèle effectif et les données de caractérisation des opérations expérimentales provient des portes à deux qubits, lesquelles ont une infidélité moyenne de 1.5 % dans l’expérience et de 0.9 % dans nos simulations. Comme l’origine exacte du bruit additionnel observé dans l’expérience n’est pas bien comprise ni quantifiée, nous évitons d’ajouter du bruit dans nos simulations afin de reproduire l’infidélité mesurée. Ce choix mène à des résultats de simulation qui présentent des performances significativement meilleures que celles observées expérimentalement, comme nous le verrons dans l’article de la section 3.4.

Il est cependant facile d’ajuster nos paramètres de simulation afin de reproduire les fidélités des portes logiques mesurées. Une approche couramment utilisée [63] consiste à ajuster le taux de déphasage des transmons lors de la simulation des opérations CZ afin d’obtenir les mêmes fidélités. Cette modification artificielle regroupe donc toutes les erreurs inexpliquées sous la forme d’un déphasage phénoménologique. L’utilisation de cette approche est illustrée à la Fig. 3.4(a) où l’on ajuste la valeur de T_φ des deux transmons afin de reproduire les données de caractérisation par étalonnage de portes aléatoires (RB) obtenues expérimentalement pour les 24 portes CZ du code de surface de distance trois.

À la Fig. 3.4(b), on observe que cette modification (symboles rouges) permet de reproduire de près la performance globale du code de surface observée dans l’expérience. En effet, on obtient un temps de cohérence logique $T_{1,L}$ de $18.2\ \mu\text{s}$ comparativement au temps de $16.4\ \mu\text{s}$ obtenu expérimentalement. Cependant, le temps de vie logique du code est une métrique globale qui dépend des nombreuses caractéristiques du dispositif et de l’algorithme de décodage utilisé. Ainsi, un nombre combinatoire de choix de paramètres de simulation différents pourrait être utilisé tout en donnant essentiellement la même performance logique. Dans la publication qui suit, nous évitons donc d’effectuer tout ajustement sur les données expérimentales et nous contentons d’utiliser les caractéristiques de cohérence et de l’Hamiltonien du système qui ont été mesurées.

3.4 Publication [1] : Réaliser la correction d’erreur quantique avec 17 transmons

L’article qui suit a été publié dans la revue *Nature* et représente la première démonstration de correction d’erreur quantique à l’aide de qubits supraconducteurs. Impliquant 17 qubits et jusqu’à 16 cycles de correction d’erreur, cette expérience nécessite l’orchestration de centaines de portes logiques à un et deux qubits et une centaine de mesures projectives, ce qui représente l’une des expériences d’informatique quantique les plus complexes réalisées à ce jour par un groupe académique. Alors que la conception, la fabrication et l’opération du dispositif ont été effectuées par nos collaborateurs à l’ETH Zurich, nous avons réalisé la modélisation et les simulations numériques de l’expérience dans le groupe à Sherbrooke. À partir d’une modélisation et d’un code initialement construits par Agustin Di Paolo, alors étudiant au doctorat à Sherbrooke, j’ai significativement étendu la pertinence et l’efficacité de nos simulations grâce à une étroite collaboration avec nos collègues à Zurich. Ainsi, j’ai pu inclure dans nos simulations l’ensemble des particularités du dispositif qui ont été caractérisées dans l’expérience, incluant les paramètres de l’Hamiltonien de chaque qubit,

leurs propriétés de cohérences et la séquence d'opération exacte utilisée en pratique. Avec ces informations, j'ai effectué toutes les simulations numériques des différentes composantes de l'expérience et ces résultats sont présentés aux Figs. 2-4 de l'article.

Nos résultats de simulation s'accordent bien avec les résultats expérimentaux au niveau de la performance totale de l'expérience quantifiée par le taux d'erreur logique p_{logique} , ainsi qu'au niveau des sous-composantes de cette expérience comme la fidélité des opérations individuelles et des stabilisateurs mesurés (lesquels impliquent 3 et 5 transmons). Étant donné que plusieurs imperfections ne sont pas incluses dans les simulations, cet accord indique que nous possédons une bonne compréhension des sources d'erreurs limitant principalement la performance de l'expérience en pratique. Les autres sources d'erreurs, lesquelles sont typiquement moins bien caractérisées, sont donc suffisamment sous contrôle afin que leur impact sur les résultats ne soit pas dominant. Sans aucun doute, avec l'amélioration du contrôle des dispositifs quantiques et la réduction des taux d'erreurs, ces imperfections devront être prises en compte et activement mitigées lorsqu'elles deviendront un facteur limitant.

Les résultats expérimentaux démontrent que le seuil de rentabilité n'a pas été atteint, lequel correspond au point où le taux d'erreur du qubit logique ($p_L \approx 3\%$) est inférieur au taux d'erreur des qubits physiques ($\sim p_{\text{phys.}}$). Nos simulations permettent de démontrer que ce seuil pourrait être atteint en diminuant les erreurs des opérations d'un facteur deux ou trois. Ces extrapolations de la performance attendue avec des taux d'erreur physiques plus bas en simulation nous permettent également de conclure que la fidélité des portes à deux qubits est le principal facteur limitant de la performance du code de surface réalisé. Les travaux de recherches permettant d'améliorer la fidélité de ces opérations à deux qubits dans le contexte de la mesure des stabilisateurs s'avèrent donc cruciaux pour démontrer une suppression des erreurs à l'aide d'une correction quantique des erreurs.

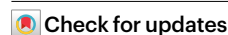
Realizing repeated quantum error correction in a distance-three surface code

<https://doi.org/10.1038/s41586-022-04566-8>

Received: 15 November 2021

Accepted: 9 February 2022

Published online: 25 May 2022



Sebastian Krinner^{1,9}✉, Nathan Lacroix^{1,9}, Ants Remm¹, Agustin Di Paolo^{2,3}, Elie Genois^{2,3}, Catherine Leroux^{2,3}, Christoph Hellings¹, Stefania Lazar¹, Francois Swiadek¹, Johannes Herrmann¹, Graham J. Norris¹, Christian Kraglund Andersen^{1,8}, Markus Müller^{4,5}, Alexandre Blais^{2,3,6}, Christopher Eichler¹ & Andreas Wallraff^{1,7}

Quantum computers hold the promise of solving computational problems that are intractable using conventional methods¹. For fault-tolerant operation, quantum computers must correct errors occurring owing to unavoidable decoherence and limited control accuracy². Here we demonstrate quantum error correction using the surface code, which is known for its exceptionally high tolerance to errors^{3–6}. Using 17 physical qubits in a superconducting circuit, we encode quantum information in a distance-three logical qubit, building on recent distance-two error-detection experiments^{7–9}. In an error-correction cycle taking only 1.1 μs , we demonstrate the preservation of four cardinal states of the logical qubit. Repeatedly executing the cycle, we measure and decode both bit-flip and phase-flip error syndromes using a minimum-weight perfect-matching algorithm in an error-model-free approach and apply corrections in post-processing. We find a low logical error probability of 3% per cycle when rejecting experimental runs in which leakage is detected. The measured characteristics of our device agree well with a numerical model. Our demonstration of repeated, fast and high-performance quantum error-correction cycles, together with recent advances in ion traps¹⁰, support our understanding that fault-tolerant quantum computation will be practically realizable.

The surface code^{4,11} is a planar realization of Kitaev's toric code³, which uses topological features of a qubit lattice to correct errors in quantum information-processing systems. This code is a prominent contender to reach fault-tolerant quantum computation because of its high error threshold of about 1% against quantum circuit noise^{5,12} and its compatibility with 2D architectures. The surface code belongs to the family of stabilizer codes^{13,14}, which encode quantum information into a joint subspace of definite parities on a set of physical data qubits to form a logical qubit. Errors are detected using measurements of auxiliary qubits to extract parity information without collapsing the logical qubit state. The fault-tolerant operation of a quantum computer requires repeated detection and correction of both bit-flip and phase-flip errors on data qubits. With an increasing number of physical qubits and thus an increasing code distance d , the number of errors $\lfloor (d-1)/2 \rfloor$ that can at least be detected and corrected per error-correction cycle increases, making the code more resilient when error rates are sufficiently low.

Error correction limited to a single type of error has been realized with repetition codes in nuclear magnetic resonance¹⁵, trapped ions¹⁶, nitrogen-vacancy centres¹⁷ and superconducting circuits^{9,18}. In single-cycle experiments, fault-tolerant stabilizer measurements and correction of both types of error have been demonstrated with the five-qubit code and the Bacon–Shor code^{19–22}. Recently, error detection in a distance-two surface code has been realized with seven qubits^{7–9},

and only very recently, repeated stabilizer-based error correction has been demonstrated with a distance-three colour code in a trapped-ion system¹⁰.

Correction of both bit-flip and phase-flip errors requires at least a distance-three code. In combination with fault-tolerant circuits for error-syndrome measurements, this guarantees that any single error on any of the constituent data and auxiliary qubits or operations can be corrected^{14,23}. While the work we discuss here focuses on digital encoding of quantum information, continuous-variable encoding, for example, in harmonic oscillator states, constitutes an alternative approach to quantum error correction; see, for example, refs.^{24–27}.

Realization in superconducting circuits

Experimentally realizing a distance-three surface code requires nine data qubits and eight auxiliary qubits^{23,28,29}, also referred to in the literature as ancilla or measurement qubits. The qubits are arranged in a diagonal, planar square lattice, the edges of which are shown in grey in the schematic of Fig. 1a. The data qubits D_j , $j = 1 \dots 9$ (red dots) form a 3×3 array and are interlaced with auxiliary qubits A_i , labelled X_i (blue) and Z_i , $i = 1 \dots 4$ (green). We realized this arrangement in a superconducting circuit using 17 transmon qubits³⁰ (yellow) capacitively coupled to each other along the edges of the square array with roughly 1-mm-long

¹Department of Physics, ETH Zurich, Zurich, Switzerland. ²Institut Quantique, Université de Sherbrooke, Sherbrooke, Québec, Canada. ³Département de Physique, Université de Sherbrooke, Sherbrooke, Québec, Canada. ⁴Institute for Quantum Information, RWTH Aachen University, Aachen, Germany. ⁵Peter Grünberg Institute, Theoretical Nanoelectronics, Forschungszentrum Jülich, Jülich, Germany. ⁶Canadian Institute for Advanced Research, Toronto, Ontario, Canada. ⁷Quantum Center, ETH Zurich, Zurich, Switzerland. ⁸Present address: QuTech and Kavli Institute for Nanoscience, Delft University of Technology, Delft, The Netherlands. ⁹These authors contributed equally: Sebastian Krinner, Nathan Lacroix. ✉e-mail: skrinner@phys.ethz.ch

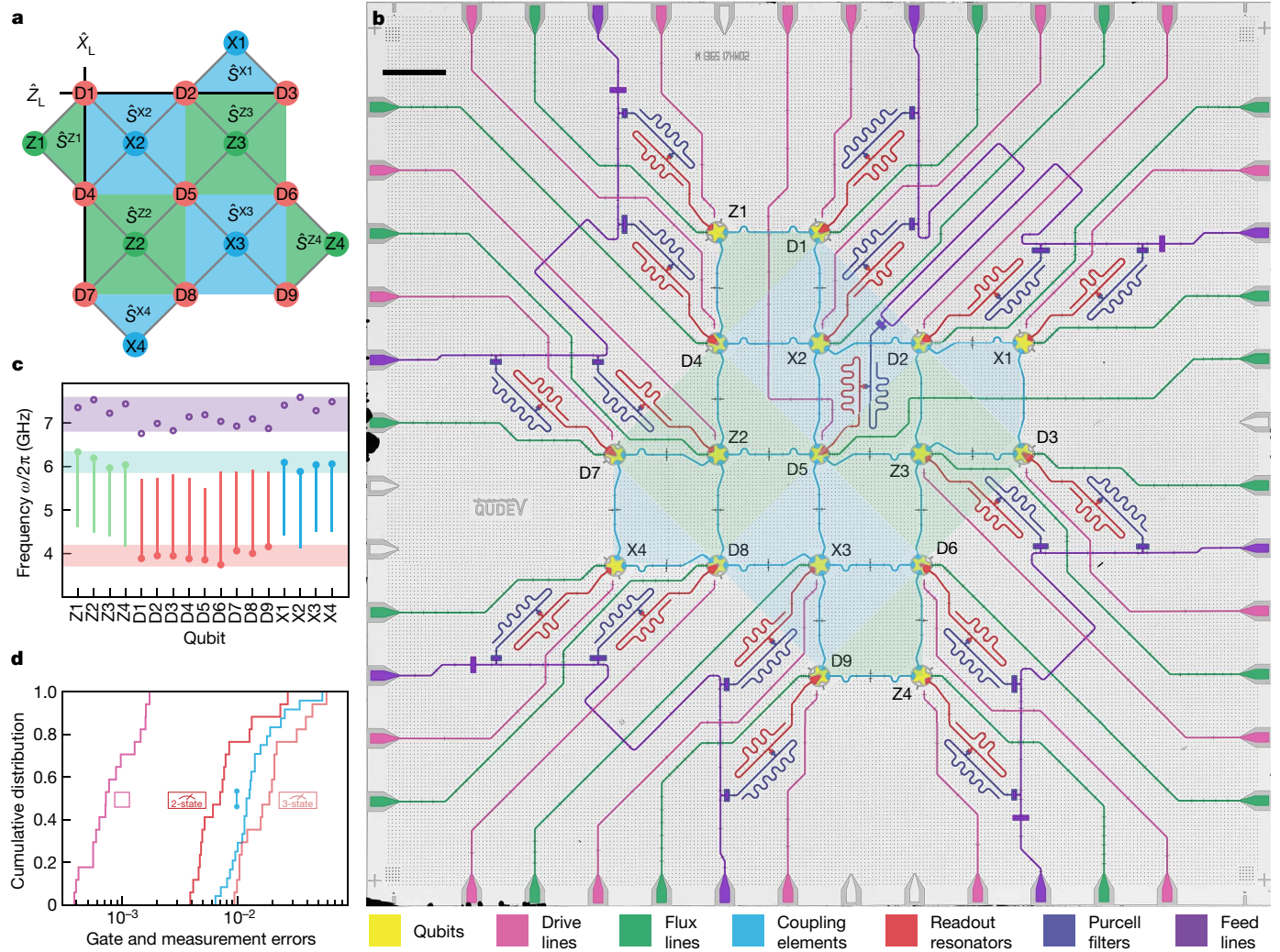


Fig. 1 | Device concept, architecture and performance. **a**, Conceptual representation of the distance-three surface code consisting of data qubits (red circles), Z-type auxiliary qubits (green circles) and X-type auxiliary qubits (blue circles), with their connectivity indicated by grey lines. The data qubits participating in the weight-three logical operators $\hat{Z}_L$ and $\hat{X}_L$ are indicated by solid black lines. Green (blue) plaquettes indicate X-type (Z-type) stabilizer circuits. **b**, False-colour micrograph of the device realizing the concept in **a** with 17 transmon qubits; see key for circuit elements and text for details. The

qubit lattice is rotated by 45° with respect to **a**. The scale bar denotes 1 mm. **c**, Frequency arrangement in three distinct bands for idling data qubits (red circles), idling Z-type/X-type auxiliary qubits (green/blue circles) and readout resonators (violet open circles). The qubit-frequency tuning ranges are indicated by vertical bars. **d**, Cumulative distributions (integrated histograms) of single-qubit gate (pink), simultaneous two-qubit gate (cyan), and two-state (red) and three-state readout errors (light red).

coplanar waveguide segments (cyan); see Fig. 1b. We discuss the fabrication of this device in the Supplementary Information (section I).

Using the auxiliary qubits X_i and Z_i , we measure the parity of the neighbouring two or four data qubits D_j , which are located at the vertices of the blue and green plaquettes (Fig. 1a, b), in the X or Z basis, see Methods. These parity measurements are also referred to as stabilizer measurements^{13,14}. The corresponding mutually commuting weight-two (or weight-four) stabilizer operators $\hat{S}^{X_i} = \prod_{j=1}^{2(4)} \hat{X}_j$ and $\hat{S}^{Z_i} = \prod_{j=1}^{2(4)} \hat{Z}_j$ of the surface code are products of two (or four) Pauli- $\hat{X}$ or Pauli- $\hat{Z}$ operators of the data qubits j located at the vertices of a given data-qubit plaquette. Measurements of the stabilizer operator $\hat{S}^{A_i}$ with outcomes $s^{A_i} = \pm 1$ indicate even or odd parity of the corresponding data-qubit state. A change of data-qubit parity signals an error. In our experiments, in which we do not reset auxiliary qubits, odd parity is indicated by a change of auxiliary-qubit state from cycle to cycle, whereas even parity is indicated by the absence of such changes.

In our experiments, the stabilizer gate sequence is realized as two or four controlled-phase (CZ) gates^{31–33} (see Methods), between data

and auxiliary qubits, operated in a low and high frequency band (see Fig. 1c), respectively, combined with initial and final $\pi/2$ rotations on the auxiliary qubits (Fig. 2a, b). The gate sequence for measuring $\hat{S}^{X_i}$ contains further initial and final $\pi/2$ rotations acting on the data qubits, implementing a basis change from the Z to the X basis (blue dashed squares in Fig. 2a, b). We apply echo pulses to the data qubits in the middle of the gate sequence to reduce dephasing of the data qubits and residual coherent coupling to spectator qubits³⁴.

The 24 pairwise CZ gates have a mean duration of 98(7) ns, including two conservatively chosen 15-ns-long buffers at the beginning and the end, and show a mean gate error of 0.015(10). The gate error histogram, shown as an integrated (cumulative) distribution, shows variations of about a factor of four in two-qubit gate error (Fig. 1d and Methods). Single-qubit gates showing a mean error of 0.0009(4) are realized by applying short resonant microwave pulses to each qubit individually through a dedicated drive line (pink coplanar waveguide in Fig. 1b). We determined both single-qubit and two-qubit gate fidelities in randomized benchmarking experiments. We discuss the experimental setup used to realize these gates in the Supplementary Information (section IV).

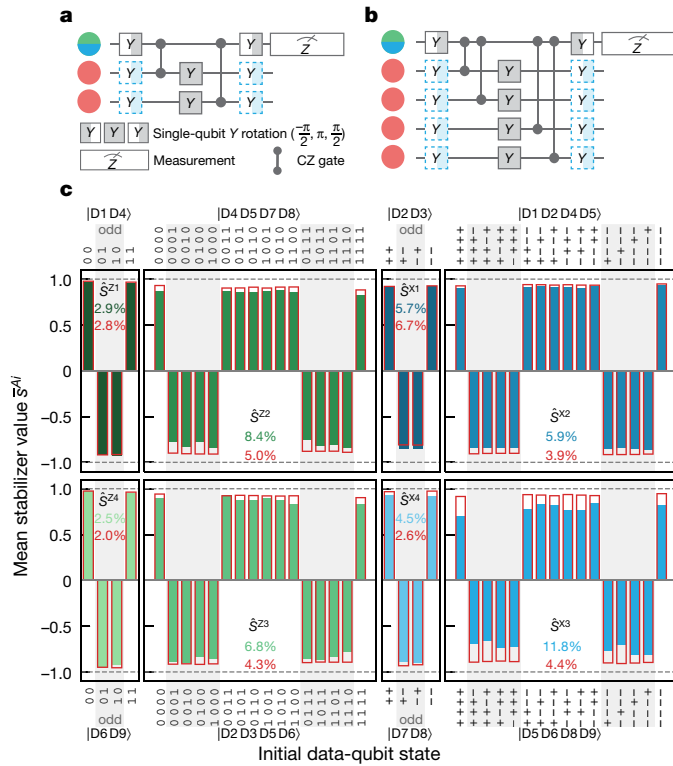


Fig. 2 | Stabilizer circuits and their characterization. Quantum circuit diagrams for weight-two (a) and weight-four (b) stabilizers acting between data qubits located on the vertices of a plaquette (red circles) and the corresponding auxiliary qubit used for Z-type (green circles) or X-type (blue circles) stabilizer measurements, in which the latter require a basis change (blue dashed squares); see text for details. c, Measured (filled bars) and simulated (red wireframes) average stabilizer values $\bar{s}^{Ai}$ versus data-qubit input state, ordered by number of excitations. Percentages are the experimental (blue and green) and simulated (red) error of $\bar{s}^{Ai}$. Grey background indicates odd parity and white indicates even parity.

A key element of individual stabilizer measurements are fast and high-fidelity measurements of auxiliary-qubit states while leaving data-qubit states unaffected^{35–37}; see Methods. Accurate stabilizer measurements using mid-cycle qubit readout on timescales comparable with or shorter than the cumulated gate times per error-correction cycle also contribute to maximizing the performance of error detection and correction in our surface-code implementation as a whole.

With our readout scheme, we discriminate the two computational qubit states and a leakage state, which—if undetected or uncorrected for—is detrimental to any surface-code implementation^{37–42}. We achieve a mean readout assignment error of 0.009(7) when discriminating the computational states only (two-state readout) and of 0.022(14) when discriminating the computational states and the leakage state (three-state readout); see Supplementary Information (section VI). The corresponding cumulative distribution for two-state readout exhibits performance variations on the device of about a factor of two when disregarding two outliers, whereas the distribution for three-state readout shows variations larger by about a factor of two (Fig. 1d).

With all elements in place for realizing a surface code, we first characterize the measurements of individual $\hat{S}^{Ai}$ stabilizers. To do so, we prepare the data qubits D_j of a given weight-two or weight-four plaquette sequentially in each one of its $2^2 = 4$ or $2^4 = 16$ basis states composed of $|0\rangle$ and $|1\rangle$ for $\hat{S}^{Zi}$ and $|+\rangle$ and $|-\rangle$ for $\hat{S}^{Xi}$. At the beginning of each experiment, all qubits are initialized by heralding the ground state $|0\rangle$ from single-shot readout.

For each input state, we compute the mean values $\bar{s}^{Ai}$ from about 4×10^4 measurements of s^{Ai} (coloured bars in Fig. 2c) and find good qualitative agreement with master equation simulations (red outlines); see Supplementary Information (sections VII and VIII) for details. Here $+1/-1$ indicate even/odd parity of the measured state. The coloured percentage values show the corresponding experimental and simulated errors of the stabilizer measurements. On average, the experimental parity assignment error is 3.9(1.3)% for weight-two stabilizers and 8.2(2.2)% for weight-four stabilizers. We attribute the differences between measurements and simulation mostly to two-qubit gate errors owing to microscopic defect modes changing their frequency on time-scales of hours or days (Supplementary Information, section III).

Having verified that all stabilizer measurements perform at high quality levels individually, we combine the stabilizer measurements into a surface-code cycle. Executing this cycle once, we prepare one of the four cardinal logical qubit states. Executing the cycle several times, we stabilize the logical states and investigate the performance of our realization of the code.

The surface-code cycle

At the beginning of each experimental sequence, we prepare the nine data qubits in either one of the product states $|0\rangle^{\otimes 9}$ and $\hat{X}_1|0\rangle^{\otimes 9}(|+\rangle^{\otimes 9}$ and $\hat{Z}_1|+\rangle^{\otimes 9}$) to begin the process of initializing the cardinal logical qubit states $|0\rangle_L$ and $|1\rangle_L$ ($|+\rangle_L$ and $|-\rangle_L$). The cardinal states are eigenstates of the eight stabilizer operators $\hat{S}^{Ai}$ and ± 1 eigenstates of the logical Pauli operators, which we choose as $\hat{Z}_L = \hat{Z}_1\hat{Z}_2\hat{Z}_3$ and $\hat{X}_L = \hat{X}_1\hat{X}_4\hat{X}_5$; see solid black lines in Fig. 1a. As required, $\hat{Z}_L$ and $\hat{X}_L$ commute with all stabilizers and anti-commute with each other. Because each of the prepared product states is an equal superposition of 16 equivalent instances of the target logical state, executing a single quantum error-correction cycle deterministically initializes the target logical state in the stabilizer eigenspace corresponding to the measurement outcome of the stabilizers (Supplementary Information, section IX).

In a single surface-code cycle, we first execute all gate operations implementing the four $\hat{S}^{Zi}$ stabilizer measurements. We realize the necessary two-qubit gates in four time steps, in each of which we execute three CZ gates simultaneously; see Methods. Parallelizing stabilizer execution is a key technical requisite for scalable quantum error correction, in particular for operation of larger-distance codes.

The gate execution is followed by readout of the Z-type auxiliary qubits to complete the Z-type stabilizer measurements; see Fig. 3a. Simultaneous with the Z-type auxiliary-qubit readout, we start executing the X-type stabilizer circuits, which are equivalent up to a basis change of the data qubits. This allows us to execute the $\hat{S}^{Zi}$ and $\hat{S}^{Xi}$ stabilizer measurements in a parallel, pipelined approach⁴³, which imposes fewer constraints on the choice of two-qubit gate interaction frequencies compared with parallel scheduling^{29,44}. See Fig. 3a for a full circuit diagram and Supplementary Information (section X) for a full pulse sequence. In the pipelined approach with fast readout, we achieve a short quantum error-correction cycle time of $t_c = 1.1 \mu\text{s}$. A short t_c relative to the physical qubit coherence times is essential to be able to identify errors unambiguously. Moreover, a short absolute value of t_c reduces the execution time of error-corrected quantum algorithms^{45,46}.

To maximize the performance of the distance-three surface code, we have designed our device with parameters minimizing leakage on data qubits. In addition, we reject all experimental runs with residual leakage events, which we detect on auxiliary qubits during the execution of each cycle and on data qubits after the last cycle, using our three-state readout; see Methods. We discuss the fraction of experimental runs retained after leakage rejection in the Supplementary Information (section XI).

To characterize the fidelity of the nine-data-qubit logical state that we initialized deterministically using our measurement-based approach, we measure the expectation values of the 2^9 Pauli operator

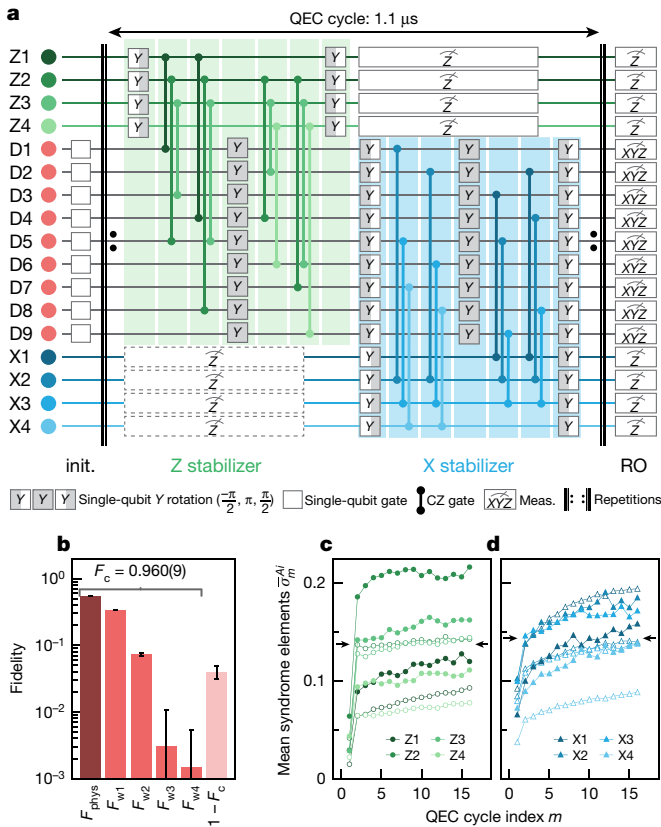


Fig. 3 | The surface-code cycle, fidelity of logical-state initialization and average error syndromes. **a**, Quantum circuit used to initialize and (repeatedly) error correct the distance-three surface code logical qubit. Green (blue) shaded circuit elements represent the parallel execution of the four Z-type (X-type) stabilizer circuits. Empty squares indicate single-qubit rotations on data qubits. In the first cycle, the X_i auxiliary qubits are not measured (dashed boxes). **b**, Fidelity of the prepared logical state $|0\rangle_L$ (dark red bar), all correctable states differing from $|0\rangle_L$ by one to four Pauli errors (red bars) and all uncorrectable states (light red bar). Fidelities are computed as the means over 500 different correctable subspaces (Supplementary Information, section IX) and error bars indicate one standard deviation. Average syndrome elements $\bar{\sigma}_m^{Ai}$ as a function of the cycle index m for the $\hat{Z}_L$ (**c**, filled circles) and the $\hat{X}_L$ (**d**, filled triangles) preservation experiment. Open symbols are simulations. The black arrows indicate the average syndrome element over all stabilizers and cycles; see text for details. QEC, quantum error correction.

strings that form the basis of the target state^{20,47}. For $|0\rangle_L$, we find a quantum state fidelity of $F_{\text{phys}} = \text{Tr}(\rho|0\rangle_L\langle 0|_L) = 54.0(1)\%$ (dark red bar in Fig. 3b). Considering only errors in the logical subspace⁷, we find a logical fidelity of $F_L = 99.6(2)\%$; see Methods and Supplementary Information (section IX). We also determine the fidelity of the experimentally prepared state with respect to correctable subspaces including states that are equivalent to the target state up to Pauli errors of weight $i = 1 \dots 4$ and find $(F_{w1}, F_{w2}, F_{w3}, F_{w4}) = [34.3(0), 7.3(5), 0.3(8), 0.1(4)]\%$. Hence the prepared initial state has a fidelity of $F_c = F_{\text{phys}} + \sum_{i=1}^4 F_{wi} = 96.0(9)\%$ with respect to states that are, in principle, correctable; see Supplementary Information (section IX).

Repeated quantum error correction

Once the first quantum error-correction cycle completes the logical-state initialization, we make use of all subsequent cycles for logical-state preservation. In our experiments, we preserve the cardinal logical qubit states for up to $n = 16$ cycles. In each cycle $m = 1 \dots n$, we extract eight stabilizer values s_m^{Ai} . Changes in stabilizer values signal

the occurrence of errors and are used to construct error syndromes σ_m consisting of eight syndrome elements $\sigma_m^{Ai} = (1 - s_m^{Ai} \times s_{m-1}^{Ai})/2$. The elements σ_m^{Ai} are inferred in each cycle from the current (m) and the previous ($m - 1$) measured stabilizer values, with $\sigma_m^{Ai} = 1(0)$ indicating an error (no error)¹⁸.

We collectively process successive syndromes σ_m to determine which data and auxiliary qubits have most probably suffered an error^{4,14,44} using the approach described in Methods. Executing n consecutive error-correction cycles and rejecting runs in which leakage has been detected, we record stabilizer measurement outcomes and construct syndromes from their values, the averages of which are shown in Fig. 3c, d. When averaging the syndromes over all individual elements and time, we find that the average syndrome element $\bar{\sigma} = 0.14 \ll 1$ is small (see arrows in Fig. 3c, d), indicating that errors are rare and therefore allowing for efficient error detection and correction⁹. In Methods, we discuss the syndrome measurement outcomes in more detail.

To determine the performance of our distance-three surface code, we extract the logical error per cycle when preserving eigenstates of the logical qubit operators $\hat{Z}_L$ and $\hat{X}_L$ versus the number of executed cycles n . For each sequence of cycles, we decode the error syndromes, including a syndrome determined from the final data-qubit readout. We use a minimum-weight perfect-matching algorithm, the weights of which we determine in an error-model-free approach (Methods). After the n th cycle, we perform a projective readout of the final data-qubit state in the Z or X basis, from which we determine the eigenvalue $z_L = \pm 1$ of $\hat{Z}_L$ or $x_L = \pm 1$ of $\hat{X}_L$.

Decoding the error syndromes and applying corrections to z_L or x_L when indicated (Methods), we compute the mean logical qubit expectation values $\bar{z}_L = \langle \hat{Z}_L \rangle$ and $\bar{x}_L = \langle \hat{X}_L \rangle$ as a function of n from a total of 10^6 experimental runs, in which the available data are reduced by ground-state heralding before and leakage rejection during each run. We observe an exponential decay of $\langle \hat{Z}_L \rangle$ and $\langle \hat{X}_L \rangle$ with n (solid symbols in Fig. 4a, b). From the logical qubit operator expectation values, we extract the logical error probability $E_L = (1 - \langle \hat{Z}_L \rangle)/2$ (solid symbols in Fig. 4c) as a function of n and find a small logical error per cycle of $\varepsilon_L = [1 - \exp(-t_c/T_{1,L})]/2 \approx t_c/2T_{1,L} = 0.032(1)$, also indicated on the right-hand side of the corresponding dataset in Fig. 4c. Equivalently, we obtain $\varepsilon_L = [1 - \exp(-t_c/T_{2,L})]/2 \approx t_c/2T_{2,L} = 0.029(1)$ for $\langle \hat{X}_L \rangle$. We note that both the coherence time of $T_{2,L} = 18.2(5) \mu\text{s}$ and the lifetime $T_{1,L} = 16.4(8) \mu\text{s}$ of the logical qubit, as extracted from the decay curves of $\langle \hat{X}_L \rangle$ and $\langle \hat{Z}_L \rangle$ are much longer than the duration of the quantum error-correction cycle $t_c = 1.1 \mu\text{s}$ in our implementation of the surface code. For completeness, in the Supplementary Information (section XI), we also state and discuss the logical error per cycle when leakage is not rejected.

We note that, in our implementation of the surface code, the logical coherence time is only about a factor of two lower than the mean physical coherence time $\bar{T}_2^* = 37.5 \mu\text{s}$ of all 17 qubits, whereas the logical relaxation time is about a factor of four lower than twice the mean physical energy relaxation time $2T_1 = 65.0 \mu\text{s}$, as indicated in Fig. 4a, b. We discuss this result further in Methods and put it into context in the next section.

Performance assessment and projection

We compare the state-preservation experiments with numerical simulations using a Monte Carlo wavefunction method (open symbols in Fig. 4a–c). The model underlying the simulations (Supplementary Information, section VII) uses the measured coherence times, interaction rates and readout errors of the device as inputs. We find that, despite the complexity of the quantum error-correction cycle, the measured expectation values $\langle \hat{X}_L \rangle$ and $\langle \hat{Z}_L \rangle$, when rejecting leakage, come close to the simulated values (open symbols in Fig. 4a, b). The logical lifetimes extracted from exponential fits (dashed lines in Fig. 4a, b) to the simulated expectation values are approximately $26 \mu\text{s}$. They

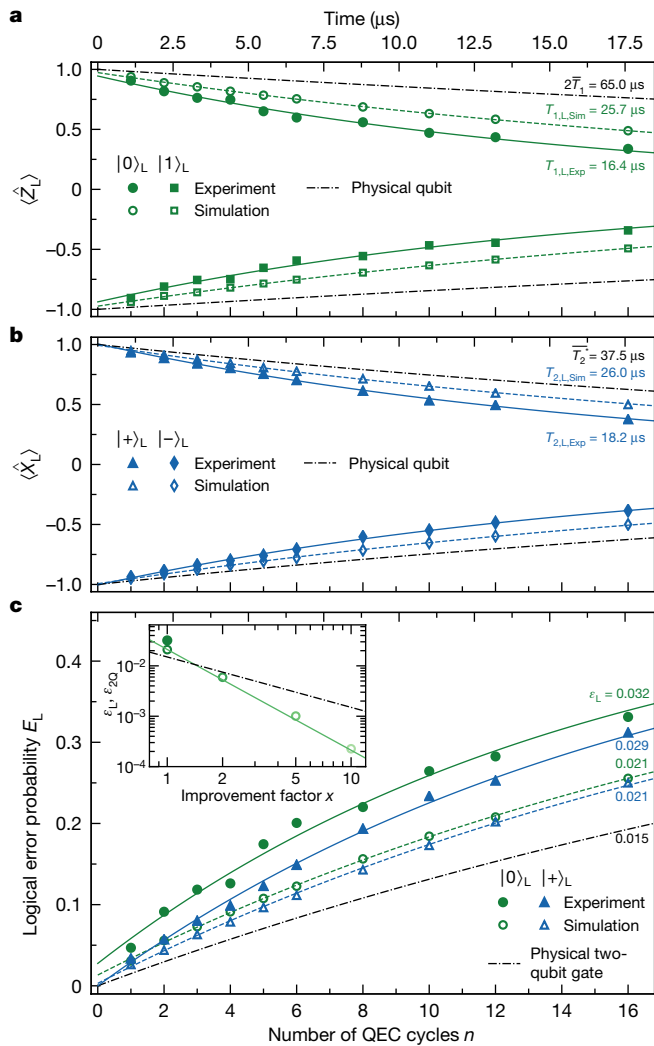


Fig. 4 | Logical-state preservation and error per cycle. **a**, Experimentally determined (full symbols) and simulated (open symbols) expectation value of the Z_L operator for prepared $|0\rangle_L$ (circles) and $|1\rangle_L$ (squares) and exponential fit (solid and dashed lines). For reference, twice the average physical qubit decay is shown (dashed-dotted line). The extracted decay times are indicated on the right. **b**, Corresponding datasets for X_L in state $|+\rangle_L$ (triangles) and $|-\rangle_L$ (diamonds). For reference, the average physical qubit coherence decay is indicated (dashed-dotted line). **c**, Logical error probability E_L as a function of the number n of error-correction cycles for $|0\rangle_L$ and $|+\rangle_L$ and the extracted error per cycle ϵ_L , indicated on the right. Same symbols as in **a** and **b**. The dashed-dotted line shows the physical two-qubit gate error accumulated over n cycles, for reference. The inset shows simulation results for ϵ_L in state $|0\rangle_L$ (green open circles) for an assumed homogeneous distribution of decoherence rates and gate and readout errors, reduced by factors of $x = 2, 5$ and 10 ; see text for details. The green solid line is a fit to $1/x^2$. QEC, quantum error correction.

provide an upper bound for the performance achievable with the specified device parameters, as further error sources such as gate control errors, population loss into microscopic defect modes and measurement-induced dephasing are not included in the simulation model.

Because the numerical simulations model the performance of our quantum device well, we use the model to project how future improvements in gate and readout fidelities are expected to reduce the logical error per cycle ϵ_L . To free the projection from device-specific spread in qubit parameters, we use the average over all 17 qubits as uniform parameters (see Supplementary Information Table 1) and find good agreement with the results obtained from the qubit-specific model.

We then uniformly reduce all physical error parameters of the numerical model by a factor x , repeat the simulations of $\langle X_L \rangle$ and $\langle Z_L \rangle$ versus n and extract the mean logical error per cycle ϵ_L as a function of the improvement factor x . We find that the simulated error per cycle scales to a good approximation as $1/x^2$ (see inset in Fig. 4c), as expected for a distance-three code⁴⁴. For reference, we plot the scaled two-qubit physical error per cycle ϵ_{2Q} in the same plot (dashed-dotted line).

A metric commonly used to assess the performance of quantum error correction compares the logical error per cycle ϵ_L to the dominant error on the physical level, typically the two-qubit gate error ϵ_{2Q} (refs. 10,48). Such a comparison is particularly relevant in architectures in which the logical two-qubit gate error is dominated by errors ϵ_L in the quantum error-correction cycles belonging to or following the logical two-qubit gate operation. The number of required cycles scales in general with d in planar architectures and in architectures allowing for transversal execution of logical two-qubit gates^{5,6,23,49,50}. The good agreement between measured (about 3%) and simulated (about 2%) logical errors ϵ_L together with their simulated quadratic scaling suggests that the break-even of per-cycle logical errors with two-qubit gate errors may be in reach for modest improvements of device performance, when using leakage detection or correction.

While a comparison with the break-even point evaluates performance for a fixed code distance d , a comparison related to the error threshold is necessary to judge how far one is from reaching the desired sub-threshold regime, in which ϵ_L decreases exponentially with d . In fact, simulations on the basis of a simplified one-parameter circuit noise model^{44,51} predict a few percent logical error per cycle at threshold, a rate that is close to the logical error per cycle observed in our experiment.

In our experiments, we demonstrate the viability of realizing quantum error correction in the surface code by detecting errors during the error-correction cycle and decoding the error syndromes and correcting for errors in post-processing, which is sufficient in a quantum memory setting. Next-generation experiments will provide the capability of correcting errors during the cycle using real-time decoding¹⁰ and fast in-sequence feedback³⁶, implemented with dedicated digital electronics. Feedback will also enable the mid-cycle suppression of leakage⁵², for example, by auxiliary qubit reset. Realizing larger surface code lattices while improving the performance of their components and demonstrating exponential suppression of logical errors with increasing code distance are upcoming important steps towards achieving the long-term goal of fault-tolerant quantum computation.

Online content

Any methods, additional references, Nature Research reporting summaries, source data, extended data, supplementary information, acknowledgements, peer review information; details of author contributions and competing interests; and statements of data and code availability are available at <https://doi.org/10.1038/s41586-022-04566-8>.

1. Preskill, J. Quantum computing in the NISQ era and beyond. *Quantum* **2**, 79 (2018).
2. Shor, P. W. Fault-tolerant quantum computation. In *Proc. 37th Conference on Foundations of Computer Science* 56 pp (IEEE, 1996).
3. Kitaev, A. Y. Fault-tolerant quantum computation by anyons. *Ann. Phys.* **303**, 2–30 (2003).
4. Dennis, E., Kitaev, A., Landahl, A. & Preskill, J. Topological quantum memory. *J. Math. Phys.* **43**, 4452–4505 (2002).
5. Raussendorf, R. & Harrington, J. Fault-tolerant quantum computation with high threshold in two dimensions. *Phys. Rev. Lett.* **98**, 190504 (2007).
6. Bombin, H. & Martin-Delgado, M. A. Quantum measurements and gates by code deformation. *J. Phys. A Math. Theor.* **42**, 095302 (2009).
7. Andersen, C. K. et al. Repeated quantum error detection in a surface code. *Nat. Phys.* **16**, 875–880 (2020).
8. Marques, J. F. et al. Logical-qubit operations in an error-detecting surface code. *Nat. Phys.* **18**, 80–86 (2022).
9. Chen, Z. et al. Exponential suppression of bit or phase errors with cyclic error correction. *Nature* **595**, 383–387 (2021).
10. Ryan-Anderson, C. et al. Realization of real-time fault-tolerant quantum error correction. *Phys. Rev. X* **11**, 041058 (2021).

11. Bravyi, S. B. & Kitaev, A. Y. Quantum codes on a lattice with boundary. Preprint at <https://arxiv.org/abs/quant-ph/9811052> (1998).
12. Wang, C., Harrington, J. & Preskill, J. Confinement-Higgs transition in a disordered gauge theory and the accuracy threshold for quantum memory. *Ann. Phys.* **303**, 31–58 (2003).
13. Gottesman, D. *Stabilizer Codes and Quantum Error Correction*. PhD thesis, California Institute of Technology (1997).
14. Terhal, B. M. Quantum error correction for quantum memories. *Rev. Mod. Phys.* **87**, 307 (2015).
15. Moussa, O., Baugh, J., Ryan, C. A. & Laflamme, R. Demonstration of sufficient control for two rounds of quantum error correction in a solid state ensemble quantum information processor. *Phys. Rev. Lett.* **107**, 160501 (2011).
16. Schindler, P. et al. Experimental repetitive quantum error correction. *Science* **332**, 1059–1061 (2011).
17. Walderherr, G. et al. Quantum error correction in a solid-state hybrid spin register. *Nature* **506**, 204–207 (2014).
18. Kelly, J. et al. State preservation by repetitive error detection in a superconducting quantum circuit. *Nature* **519**, 66–69 (2015).
19. Knill, E., Laflamme, R., Martinez, R. & Negrevergne, C. Benchmarking quantum computers: the five-qubit error correcting code. *Phys. Rev. Lett.* **86**, 5811 (2001).
20. Abobeih, M. H. et al. Fault-tolerant operation of a logical qubit in a diamond quantum processor. Preprint at <https://arxiv.org/abs/2108.01646> (2021).
21. Egan, L. et al. Fault-tolerant control of an error-corrected qubit. *Nature* **598**, 281–286 (2021).
22. Hilder, J. et al. Fault-tolerant parity readout on a shuttling-based trapped-ion quantum computer. *Phys. Rev. X* **12**, 011032 (2022).
23. Horsman, C., Fowler, A. G., Devitt, S. & Meter, R. V. Surface code quantum computing by lattice surgery. *New J. Phys.* **14**, 123011 (2012).
24. Ofek, N. et al. Extending the lifetime of a quantum bit with error correction in superconducting circuits. *Nature* **536**, 441–445 (2016).
25. Hu, L. et al. Quantum error correction and universal gate set operation on a binomial bosonic logical qubit. *Nat. Phys.* **15**, 503–508 (2019).
26. Flühmann, C. et al. Encoding a qubit in a trapped-ion mechanical oscillator. *Nature* **566**, 513–517 (2019).
27. Campagne-Ibarcq, P. et al. Quantum error correction of a qubit encoded in grid states of an oscillator. *Nature* **584**, 368–372 (2020).
28. Bombin, H. & Martin-Delgado, M. A. Optimal resources for topological two-dimensional stabilizer codes: comparative study. *Phys. Rev. A* **76**, 012305 (2007).
29. Tomita, Y. & Svore, K. M. Low-distance surface codes under realistic quantum noise. *Phys. Rev. A* **90**, 062320 (2014).
30. Koch, J. et al. Charge-insensitive qubit design derived from the Cooper pair box. *Phys. Rev. A* **76**, 042319 (2007).
31. Strauch, F. W. et al. Quantum logic gates for coupled superconducting phase qubits. *Phys. Rev. Lett.* **91**, 167005 (2003).
32. DiCarlo, L. et al. Preparation and measurement of three-qubit entanglement in a superconducting circuit. *Nature* **467**, 574–578 (2010).
33. Negirneac, V. et al. High-fidelity controlled-Z gate with maximal intermediate leakage operating at the speed limit in a superconducting quantum processor. *Phys. Rev. Lett.* **126**, 220502 (2021).
34. Krinner, S. et al. Benchmarking coherent errors in controlled-phase gates due to spectator qubits. *Phys. Rev. Appl.* **14**, 024042 (2020).
35. Negnevitsky, V. et al. Repeated multi-qubit readout and feedback with a mixed-species trapped-ion register. *Nature* **563**, 527–531 (2018).
36. Andersen, C. K. et al. Entanglement stabilization using ancilla-based parity detection and real-time feedback in superconducting circuits. *npj Quantum Inf.* **5**, 69 (2019).
37. Bultink, C. C. et al. Protecting quantum entanglement from leakage and qubit errors via repetitive parity measurements. *Sci. Adv.* **6**, eaay3050 (2020).
38. Aliferis, P. & Terhal, B. M. Fault-tolerant quantum computation for local leakage faults. *Quantum Inf. Comput.* **7**, 139–156 (2007).
39. Fowler, A. G. Coping with qubit leakage in topological codes. *Phys. Rev. A* **88**, 042308 (2013).
40. Ghosh, J. & Fowler, A. G. Leakage-resilient approach to fault-tolerant quantum computing with superconducting elements. *Phys. Rev. A* **91**, 020302 (2015).
41. Suchara, M., Cross, A. W. & Gambetta, J. M. Leakage suppression in the toric code. *Quantum Inf. Comput.* **15**, 997–1016 (2015).
42. Varbanov, B. M. et al. Leakage detection for a transmon-based surface code. *npj Quantum Inf.* **6**, 102 (2020).
43. Versluis, R. et al. Scalable quantum circuit and control for a superconducting surface code. *Phys. Rev. Appl.* **8**, 034021 (2017).
44. Fowler, A. G., Mariantoni, M., Martinis, J. M. & Cleland, A. N. Surface codes: towards practical large-scale quantum computation. *Phys. Rev. A* **86**, 032324 (2012).
45. von Burg, V. et al. Quantum computing enhanced computational catalysis. *Phys. Rev. Res.* **3**, 033055 (2021).
46. Babbush, R. et al. Focus beyond quadratic speedups for error-corrected quantum advantage. *PRX Quantum* **2**, 010103 (2021).
47. Nigg, D. et al. Quantum computations on a topologically encoded qubit. *Science* **345**, 302–305 (2014).
48. Trout, C. J. et al. Simulating the performance of a distance-3 surface code in a linear ion trap. *New J. Phys.* **20**, 043038 (2018).
49. Landahl, A. J. & Ryan-Anderson, C. Quantum computing by color-code lattice surgery. Preprint at <https://arxiv.org/abs/1407.5103> (2014).
50. Gutiérrez, M., Müller, M. & Bermúdez, A. Transversality and lattice surgery: exploring realistic routes toward coupled logical qubits with trapped-ion quantum processors. *Phys. Rev. A* **99**, 022330 (2019).
51. Stephens, A. M. Fault-tolerant thresholds for quantum error correction with the surface code. *Phys. Rev. A* **89**, 022321 (2014).
52. McEwen, M. et al. Removing leakage-induced correlated errors in superconducting quantum error correction. *Nat. Commun.* **12**, 1761 (2021).

Publisher's note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

© The Author(s), under exclusive licence to Springer Nature Limited 2022

Methods

Parity measurement

In a Z-basis measurement, if an odd number of the data qubits involved in the parity operator under consideration is in the $|1\rangle$ state, the auxiliary qubit state is flipped. On the other hand, if an even number of data qubits is in $|1\rangle$, the auxiliary qubit state remains unchanged. The equivalent is true in the X basis for an even or an odd number of data qubits in the $|-\rangle$ state. Here $|0\rangle$ and $|1\rangle$ are the transmon qubit ground and first excited states, respectively, and $|\pm\rangle = (|0\rangle \pm |1\rangle)/\sqrt{2}$ are their superpositions. To map the parity of the data qubits D_j onto the corresponding auxiliary qubit, we effectively use a sequence of controlled-NOT gates with the data qubits as control and the auxiliary qubit as target, and subsequently measure the state of the auxiliary qubit in the Z basis using single-shot readout. In the Z basis, a single bit-flip error of any individual data qubit leads to a change of parity, as does a single phase-flip in the X basis. Hence measurements of changes of data-qubit parities allow us to detect and identify phase-flip or bit-flip errors, as long as they occur sufficiently rarely⁴⁴.

CZ gates

We realize the necessary two-qubit CZ gates by tuning adjacent pairs of data (D_j) and auxiliary qubits (X_i, Z_i) into resonance^{31–33} using individual flux lines implemented with coplanar waveguides (green in Fig. 1b) shorted near the superconducting quantum interference device (SQUID) loop of each qubit. We fabricated all qubits with asymmetric SQUIDs^{53,54} to allow for data qubits to idle at their minimum and auxiliary qubits at their maximum frequencies, at which the qubits are first-order insensitive to flux noise. Data qubits are designed with idle frequencies 3.7–4.1 GHz in a low-frequency band and auxiliary qubits with idle frequencies 5.9–6.3 GHz in a high-frequency band (red and blue/green dots in Fig. 1c); also see Supplementary Information (section II).

We implement CZ gates by tuning both data and auxiliary qubits to an intermediate interaction frequency ω_{int} and $\omega_{\text{int}} - \alpha$, respectively, with $\omega_{\text{int}}/2\pi$ ranging from 4.4 to 5.6 GHz (Supplementary Information, section III). The qubit anharmonicity $\alpha \approx -0.17$ GHz is designed to be small to minimize residual qubit–qubit interactions³⁴. We make use of net-zero flux pulses³³, which reduce both the detrimental effect of low-frequency flux noise on qubit coherence and the impact of non-idealities in the transfer function of the flux lines on gate fidelities. Given the large designed detuning of about 2 GHz between the data and auxiliary qubits at their idle frequencies, we calculate residual ZZ interaction strengths between qubits lower than $\alpha_{zz}/2\pi \approx 8$ kHz (ref.³⁴). It is only during two-qubit gate execution that α_{zz} increases by a factor of approximately 2 to 25, depending on the interaction frequency ω_{int} , which we partially mitigate using echo pulses. The coupling strength between auxiliary qubits and data qubits at the interaction point is about $J/2\pi \approx 7$ MHz.

Two-qubit gate error

We determine the two-qubit gate error from interleaved randomized benchmarking experiments with sets of three gates executed in parallel, as used in our realization of the surface-code cycle. Time-varying microscopic defects in our device have a detrimental influence on two-qubit gate performance and are responsible for outliers in the gate-error distribution (Supplementary Information, section III).

Qubit readout

Each qubit is coupled to a resonant pair of readout resonator and Purcell filter (red and blue $\lambda/4$ coplanar waveguide resonators in Fig. 1b). Moreover, each readout resonator is coupled strongly to the qubit ($g/2\pi \approx 169$ MHz for auxiliary qubits and $g/2\pi \approx 252$ MHz for data qubits) and has a large effective bandwidth ($\kappa_{\text{eff}}/2\pi \approx 10$ MHz) to enable fast, high-fidelity readout⁵⁵. The individual Purcell filters both maintain

high qubit coherence, despite the large coupling and bandwidth of the readout resonators, and reduce undesired readout crosstalk (Supplementary Information, section V) between qubits that are in close proximity or have similar frequencies⁵⁶. This is particularly important for the simultaneous frequency-multiplexed readout of groups of four or five qubits using joint feed lines (purple coplanar waveguides in Fig. 1b). The readout resonator frequencies are separated by about 200 MHz within each feed line and occupy a frequency band extending from 6.8 to 7.6 GHz (purple points in Fig. 1c).

We read out the states of all qubits dispersively by applying frequency-multiplexed, Gaussian-filtered microwave pulses of duration 200–300 ns to all four feed lines. We integrate the transmitted signals in a heterodyne detection scheme for a duration of 400 ns (Supplementary Information, section VI). Auxiliary qubits are read out near their idle frequencies, whereas data qubits are read out at a flux-tuned qubit frequency of approximately 5 GHz, reducing the data-qubit-readout resonator detuning⁵⁷ and thus enhancing the dispersive coupling and the readout fidelity^{55,58}.

Gate sequence

The two-qubit gates are accompanied by a set of single-qubit gates applied to all auxiliary qubits in a leading and a trailing time step, and a dynamical decoupling pulse applied to all data qubits at a central time step (Fig. 3a). We choose the order of gate operations to provide resilience against single auxiliary qubit errors and to avoid interactions with microscopic defect modes (Supplementary Information, section III).

Leakage detection

We make use of a leakage-detection scheme on the basis of three-state readout, which allows us—in post-processing—to reject those sequences in which any of the qubits were measured in a leakage state; see Supplementary Information (section VI). In our CZ gate scheme, we make use of the second excited state $|2\rangle$ of the auxiliary qubits rather than that of the data qubits to mediate the interaction, which minimizes data qubit leakage. Performing three-state readout of the data qubits after the final error-correction cycle, we reject experimental runs for which data qubit leakage was detected. The rejected fraction per qubit and per cycle amounts to 0.0017(2). In addition, there is a cycle-independent rejection probability of about 0.01 per qubit, owing to false positives caused by readout error. In addition, we detect if any of the eight auxiliary qubits has leaked to the $|2\rangle$ state in any of the n cycles, using the same three-state readout, and find an average rejection probability of 0.0094(4) per qubit per cycle. In total, this leads to a rejected data fraction per cycle of 8%.

Logical-state characterization

For the data shown in Fig. 3b, we use our leakage-detection scheme and correct for readout errors on data qubits. The logical fidelity is calculated as $F_L = F_{\text{phys}}/P_L = 99.6(2)\%$, in which $P_L = 54.2(1)\%$ is the experimentally measured probability of preparing a state in the logical subspace (Supplementary Information, section IX). Both F_{phys} and P_L are smaller than in a distance-two surface code (see ref.⁷ for an example) because the two quantities are expected to decrease with increasing distance d at a constant physical error rate. To further evaluate the performance of our logical-state initialization, we analyse the fidelity of the prepared state with respect to subspaces of states that our surface-code implementation can, in principle, correct. The errors that are correctable by the distance-three surface code include all single-qubit Pauli (weight-one) errors $\hat{X}_j, \hat{Y}_j$ or $\hat{Z}_j$ on any data qubit j and a subset of higher-weight errors; see Supplementary Information (section IX) for details. We observe that weight-one errors account for most of the errors on data qubits in the $|0\rangle_L$ state initialization, with higher-weight errors having a largely reduced probability of occurrence (see Fig. 3b).

Syndrome graph

For syndrome analysis, we construct a graph in which the syndrome elements are shown at the auxiliary qubit locations along two spatial coordinates for each cycle index m , which forms the temporal coordinate. Spatial and temporal correlations between non-zero syndrome elements correspond to data and auxiliary qubit errors, respectively. If the overall error rate is sufficiently low, we obtain a low density of non-zero syndrome elements, or – equivalently – mean syndrome element values $\bar{\sigma}_m^{Ai} \ll 1$, with the mean taken over experimental realizations. In that case, the underlying errors can be decoded with low ambiguity (Supplementary Information, section XII).

Syndrome analysis

All syndrome elements $\bar{\sigma}_m^{Ai}$ averaged over repetitions of the experiments are approximately constant as a function of m for $m \geq 2$; see Fig. 3c, d obtained for preserving $|0\rangle_L$, $|1\rangle_L$ and $|+\rangle_L$, $|-\rangle_L$, respectively. We attribute the small remaining increase of $\bar{\sigma}_m^{Ai}$ with m to the fact that, in the absence of auxiliary qubit reset, auxiliary qubits initially prepared in the ground state tend to an asymptotic probability of 0.5 to be in the excited state after m cycles. As a result, with increasing m , auxiliary qubits suffer from larger decoherence during readout and during the subsequent idling periods of about 150 ns before the start of the next quantum-error-correction cycle. Our numerical simulations show the same feature (open symbols in Fig. 3c, d). For $m = 1$, the averaged syndrome elements $\bar{\sigma}_1^{Ai}$ are reduced because the corresponding reference stabilizer values are computed from the initial data qubit product state, which we prepare with high fidelity. In the first cycle, the four values of $\bar{\sigma}_1^{Zi}$ are smaller than $\bar{\sigma}_1^{Xi}$ because a quantum-error-correction cycle starts with measurements of $\hat{S}^{Zi}$ and errors thus accumulate only during half a cycle.

Error decoding

We decode the error syndromes using a minimum-weight perfect-matching algorithm^{59,60}. We determine the weights in an error-model-free approach by inferring the errors per cycle from the measured data using a correlation analysis of the syndromes as described in the Supplementary Information (section XII) and refs.^{9,61}. The correction of an error, initiated by analysing all cycles in post-processing, takes the form of changing the sign of the logical qubit operator values z_L and x_L when indicated by the decoder. We note that, for correcting $\hat{Z}_L$, it is sufficient to decode only syndromes $\{\sigma_m^{Zi}\}$, or – equivalently – only $\{\sigma_m^{Xi}\}$ for $\hat{X}_L$.

Logical coherence times

The reference for the logical energy relaxation time T_{1L} is twice the mean physical qubit energy relaxation time T_1 , assuming an infinite lifetime of the physical qubit ground state for the comparison.

We also note that, within error bars, the lifetimes of the states $|0\rangle_L$ and $|1\rangle_L$, and the coherence times of the states $|+\rangle_L$ and $|-\rangle_L$ are, respectively, identical. This is expected, as all cardinal states of the logical qubit have the same number of qubits in the excited state and as our dynamical decoupling scheme alternates between the states $|0\rangle_L$ and $|1\rangle_L$, and between the states $|+\rangle_L$ and $|-\rangle_L$, respectively, which are therefore affected similarly by decoherence.

To compare the performance of the error-correction experiments presented here to the recent error-detection experiments in distance-two surface codes realized with seven qubits^{7–9}, we post-select data from those runs of our distance-three experiment in which all syndrome elements are zero. When doing so, we estimate logical lifetimes in excess of 1 ms (Supplementary Information, section XIII). In

the seven-qubit device, logical lifetimes and coherence times in the range 60–70 μ s were observed, with comparable data-retention rates of about 50% per cycle when post-selecting on no detected errors by evaluating three stabilizers⁷. The current 17-qubit device achieves similar data-retention rates when evaluating eight stabilizers per cycle. This observation demonstrates that our 17-qubit device, including the tune-up and calibration procedures used, performs notably better than our previous seven-qubit device.

Data availability

All data are available from the corresponding author on reasonable request.

53. Strand, J. D. et al. First-order sideband transitions with flux-driven asymmetric transmon qubits. *Phys. Rev. B* **87**, 220505 (2013).
54. Hutchings, M. D. et al. Tunable superconducting qubits with flux-independent coherence. *Phys. Rev. Appl.* **8**, 044003 (2017).
55. Walter, T. et al. Rapid high-fidelity single-shot dispersive readout of superconducting qubits. *Phys. Rev. Appl.* **7**, 054020 (2017).
56. Heinsoo, J. et al. Rapid high-fidelity multiplexed readout of superconducting qubits. *Phys. Rev. Appl.* **10**, 034040 (2018).
57. Sank, D. et al. Measurement-induced state transitions in a superconducting qubit: beyond the rotating wave approximation. *Phys. Rev. Lett.* **117**, 190503 (2016).
58. Wallraff, A. et al. Approaching unit visibility for control of a superconducting qubit with dispersive readout. *Phys. Rev. Lett.* **95**, 060501 (2005).
59. O'Brien, T. E., Tarasinski, B. & DiCarlo, L. Density-matrix simulation of small surface codes under current and projected experimental noise. *npj Quantum Inf.* **3**, 39 (2017).
60. Edmonds, J. Paths, trees, and flowers. *Can. J. Math.* **17**, 449–467 (1965).
61. Spitz, S. T., Tarasinski, B., Beenakker, C. W. J. & O'Brien, T. E. Adaptive weight estimator for quantum error correction in a time-dependent environment. *Adv. Quantum Technol.* **1**, 1800012 (2018).

Acknowledgements We are grateful for valuable discussions with Q. Ficheux and C. Lledó. We acknowledge the contributions of R. Boell to the experimental setup and M. Kerschbaum for early work on the two-qubit gate implementation. The team in Zurich acknowledges financial support by the Office of the Director of National Intelligence (ODNI), Intelligence Advanced Research Projects Activity (IARPA), through the U.S. Army Research Office grant W911NF-16-1-0071, by the EU Flagship on Quantum Technology H2020-FETFLAG-2018-03 project 820363 OpenSuperQ, by the National Center of Competence in Research ‘Quantum Science and Technology’ (NCCR QSIT), a research instrument of the Swiss National Science Foundation (SNSF, grant number 51NF40-185902), by the SNSF R’Equip grant 206021-170731, by the EU programme H2020-FETOPEN project 828826 Quormorphic and by ETH Zurich. S.K. acknowledges financial support from Fondation Jean-Jacques et Félícia Lopez-Loreta and the ETH Zurich Foundation. The work in Sherbrooke was undertaken thanks in part to funding from NSERC, Canada First Research Excellence Fund, ARO grant W911NF-18-1-0411, the Ministère de l’Économie et de l’Innovation du Québec and U.S. Department of Energy, Office of Science, National Quantum Information Science Research Centers, Quantum Systems Accelerator. M.M. acknowledges support by the U.S. Army Research Office grant W911NF-16-1-0070. The views and conclusions contained herein are those of the authors and should not be interpreted as necessarily representing the official policies or endorsements, either expressed or implied, of the ODNI, IARPA or the U.S. Government.

Author contributions S.K., N.L. and A.R. planned the experiments, S.K. and N.L. performed the main experiment and S.K. and N.L. analysed the data. F.S., A.R. and C.K.A. designed the device and S.K., A.R. and G.J.N. fabricated the device. N.L., C.H. and S.L. developed the experimental software framework and A.R., C.H., N.L., S.K. and S.L. developed the control and calibration software routines. A.R., J.H., S.K. and C.H. designed and built elements of the room-temperature setup and S.K., A.R., C.H., S.L., N.L. and F.S. maintained the experimental setup. S.K., N.L., A.R., C.H., S.L. and C.K.A. characterized and calibrated the device and the experimental setup. E.G., A.D.P. and C.L. performed the numerical simulations. M.M. provided guidance on logical qubit evaluation methodology aspects. S.K., N.L., A.R., C.H. and S.L. prepared the figures for the manuscript and S.K., N.L., A.R., C.E. and A.W. wrote the manuscript, with inputs from all co-authors. A.B., C.E. and A.W. supervised the work.

Competing interests The authors declare no competing interests.

Additional information

Supplementary information The online version contains supplementary material available at <https://doi.org/10.1038/s41586-022-04566-8>.

Correspondence and requests for materials should be addressed to Sebastian Krinner.

Peer review information Nature thanks Kyungjoo Noh and the other, anonymous, reviewers for their contribution to the peer review of this work.

Reprints and permissions information is available at <http://www.nature.com/reprints>.

3.5 Perspectives futures

Outre le besoin d'améliorer essentiellement toutes les opérations réalisées dans le code de surface, les prochaines étapes vers la conception d'un ordinateur quantique tolérant aux fautes avec une telle architecture consistent à démontrer des opérations logiques entre plusieurs qubits encodés dans des codes de surface ainsi que l'opération d'un qubit logique avec un temps de cohérence macroscopique, par exemple de l'ordre de plusieurs minutes. Ces démonstrations nécessiteront l'utilisation de beaucoup plus que 17 transmons, ce qui nous oblige à utiliser des approches de simulation qui sont en mesure de décrire de tels espaces d'Hilbert afin de bien comprendre les performances attendues et éventuellement observées en pratique.

Il existe donc un besoin important de développer des outils de simulation qui sont adaptés pour des systèmes de centaines de qubits, mais qui capturent tout de même les processus physiques importants du dispositif. Un exemple important d'un tel outil est l'utilisation d'une approche semi-classique afin de décrire la fuite de niveau des transmons à l'aide de processus incohérents, une approche pouvant être mise à l'échelle pour de nombreux qubits [74]. Il serait également intéressant d'étudier les interfaces entre différents modèles de simulation permettant de conserver l'information importante à la performance précise du code tout en simplifiant le modèle, par exemple en démontrant un passage fidèle d'une simulation de plusieurs équations maîtresses (possible pour $n \lesssim 20$) à une représentation d'opérateurs de Kraus (possible pour $n \lesssim 40$) puis à une simulation de portes de Clifford (possible pour $n \lesssim 10000$).

D'un autre côté, la simulation précise des dynamiques temporelles d'opérations impliquant seulement quelques qubits demeure très utile pour améliorer notre capacité à contrôler le dispositif quantique. Nous en verrons un exemple important au chapitre suivant en démontrant la forte corrélation entre la précision du modèle de simulation utilisé et notre capacité à réaliser du contrôle quantique d'une grande fidélité. Il serait intéressant de concevoir des opérations dont les caractéristiques sont spécifiquement conçues pour augmenter la performance logique d'un code correcteur. Il est par exemple possible de démontrer que la phase acquise par un transmon lors d'une porte CZ avec un second qubit dans l'état $|2\rangle$ (c'est-à-dire un état ayant fui le sous-espace de calcul) peut avoir un impact significatif sur le taux d'erreur logique du code de surface [75, 74]. Ce degré de liberté est habituellement absent de la définition usuelle de la fidélité de porte moyenne du CZ, une démonstration que des degrés de liberté additionnels pourraient être utilisés afin d'optimiser la performance du code de correction d'erreur.

Dans l’objectif de favoriser l’exploration de nouveaux protocoles de contrôle de ce genre, je travaille actuellement à développer une librairie numérique facilitant la simulation précise des dynamiques de systèmes supraconducteurs. Cette librairie se nomme *Blueprint* et permet de combiner efficacement les composantes supraconductrices utilisées telles que les résonateurs LC, les transmons et leurs coupleurs afin de créer l’Hamiltonien et les opérateurs de dissipation nécessaires à la simulation de toutes les opérations du code de surface (portes logiques à un et deux qubits, mesure dispersive, opérations de réduction de la fuite de niveaux [76, 77], etc.). Le code utilise la librairie JAX [78], laquelle permet l’accélération des opérations réalisées sur des tenseurs (entre autres en utilisant des processeurs graphiques) et le calcul des gradients numériques par autodifférentiation. *Blueprint* est également conçue pour fonctionner directement avec la librairie *dynamics* [79] permettant d’accélérer la simulation de dynamiques quantiques produites par exemple par une équation maîtresse. Je présenterai la librairie *dynamics* à la section 6.4. Comme nous le verrons au chapitre suivant, ces caractéristiques permettent de réaliser diverses optimisations par descente de gradient sur le système quantique, ouvrant la porte à de nombreuses applications en contrôle optimal quantique et en estimation de paramètres. La librairie *Blueprint* a le potentiel de grandement simplifier et d’accélérer le genre de simulations réalisées dans ce chapitre, lesquelles ont été construites à partir de zéro au niveau du modèle physique et qui ont plutôt recours à la librairie *QuTiP* [68] pour résoudre l’équation maîtresse.

Chapitre 4

Contrôle optimal basé sur l'apprentissage automatique

Dans ce chapitre, j'introduis une approche permettant de réaliser des opérations quantiques plus rapides et de plus grandes fidélités sur un dispositif quantique donné. Cette méthode est conçue pour être facilement utilisable en pratique. Plus spécifiquement, j'adopte une approche d'apprentissage automatique afin de construire un modèle des dynamiques quantiques directement à partir de données obtenues en mesurant le système. Par la suite, j'utilise ce modèle entraîné dans une boucle d'optimisation permettant de trouver les contrôles qui réalisent les opérations voulues avec une fidélité optimale. Je démontre l'utilité de cette approche à l'aide de simulations numériques des dynamiques d'un transmon.

Dans ce qui suit, j'introduis d'abord une vue d'ensemble de la problématique de recherche ainsi que sur les avantages et limitations des solutions existantes. Je présente ensuite les concepts clés couverts par notre approche à la section 4.2, tels que le contrôle optimal quantique, l'apprentissage automatique, ainsi que leur union pour le contrôle de dispositifs quantiques. Je résume les principaux résultats de ce projet de recherche à la section 4.3. La description complète de notre approche ainsi que l'ensemble des résultats se retrouvent dans l'article présenté à la section 4.4. Finalement, je présente diverses avenues de recherche et applications associées à la méthodologie que j'ai développée à la section 4.5, lesquelles permettent de s'attaquer à des problèmes concrets dans le développement de technologies quantiques.

4.1 Mise en contexte

Notre capacité à contrôler avec précision l'information quantique contenue dans un système est étroitement liée à notre compréhension des dynamiques en jeu dans ce système. Étant donné la complexité de ces dynamiques pour un système quantique avec de nombreuses composantes et un environnement subissant des fluctuations temporelles, il est très difficile d'obtenir un modèle précis de tous ces degrés de liberté. Dans l'objectif de réaliser des opérations quantiques de haute fidélité, il est alors nécessaire de caractériser spécifiquement les processus de bruit dominants et d'ajuster notre conception et nos contrôles du dispositif afin d'atténuer leurs effets néfastes. Considérons deux exemples d'une telle approche pour le transmon.

Une source majeure de décohérence pour les transmons provient des pertes diélectriques associées à la présence de systèmes à deux niveaux situés dans les matériaux entourant les qubits [80]. La présence de ces modes non contrôlés est principalement due à des imperfections de fabrication menant à l'existence de défauts dans le réseau atomique des matériaux. Ces défauts se trouvent notamment aux interfaces entre le substrat sur lequel le qubit est imprimé et le métal formant le qubit supraconducteur, à la surface du diélectrique formant la jonction Josephson, ainsi qu'aux interfaces substrat-air, métal-air et métal-métal [81]. Ces imperfections sont d'ailleurs à l'origine de fluctuations temporelles importantes du temps de relaxation des dispositifs supraconducteurs [82, 83].

Alors que les processus de fabrication font l'objet de recherche active afin de limiter la présence de ces systèmes à deux niveaux [84, 85, 86, 87], une caractérisation précise de leur impact sur la cohérence des qubits supraconducteurs permet de grandement restreindre leur impact négatif sur la cohérence et la fidélité des opérations quantiques. En effet, il est possible d'optimiser, et ce jusqu'à grande échelle, les fréquences utilisées lors de toutes les opérations des qubits à fréquence variable afin d'éviter les régions où le couplage avec un système à deux niveaux est important [88, 89]. Cette allocation judicieuse des fréquences utilise le fait que le couplage est inversement proportionnel au décalage de fréquence entre le système à deux niveaux et le qubit. Cependant, ce processus d'optimisation doit être répété périodiquement, typiquement en ajustant les fréquences tous les jours à la suite du mouvement dans le domaine fréquentiel des défauts. Une approche alternative fonctionnant avec des qubits à fréquence fixe consiste à modifier la fréquence même des systèmes à deux niveaux en appliquant des contrôles externes. Par exemple, il a été démontré que l'application de signaux électriques, de déformations physiques au substrat, et d'impulsions micro-ondes peuvent modifier le spectre fréquentiel des défauts et du même coup modifier leur impact sur la cohérence du système quantique étudié [90, 91, 92, 93].

Un second exemple proéminent de l'importance d'une caractérisation adéquate dans la réalisation d'opérations quantiques de haute fidélité provient de la distorsion des impulsions électromagnétiques que l'on envoie au dispositif. Ces distorsions comprennent tout écart entre les impulsions qui sont générées dans les appareils électroniques de contrôle et les impulsions réellement reçues par le qubit. Ces distorsions sont dues entre autres aux pertes lors de la propagation des signaux dans les lignes de transmission, aux réflexions aux diverses interfaces où l'impédance des milieux de propagation n'est pas parfaitement accordée, et aux imperfections des composantes micro-ondes utilisées telles que le mélangeur IQ et les filtres micro-ondes. La déformation résultant de l'ensemble de ces effets est typiquement modélisée par une fonction de transfert, une simple transformation linéaire. Afin de réaliser des opérations quantiques de haute fidélité, il est alors nécessaire de caractériser cette transformation et de l'inverser numériquement afin de générer des signaux qui, une fois propagés jusqu'au dispositif quantique, ont la forme désirée. Cette caractérisation est typiquement réalisée à température pièce avec des appareils d'analyse de signaux standards tels qu'un oscilloscope ou un analyseur de réseau vectoriel (VNA) [22, 94]. Cependant, les propriétés de transmission changent avec la température et les transmons se retrouvent à des températures cryogéniques. De plus, cette caractérisation n'inclut pas l'ensemble de la propagation jusqu'aux qubits montés sur leur porte-échantillon, où les différents contacts et interfaces peuvent mener à des réflexions importantes.

À mesure que les techniques de fabrication et que la cohérence des qubits supraconducteurs s'améliorent, la contribution des distorsions dans l'erreur totale des opérations quantiques devient de plus en plus grande. Afin de réaliser des opérations sous le seuil d'erreur requis par la correction d'erreur quantique, il est alors nécessaire d'améliorer notre caractérisation et notre contrôle des signaux que l'on envoie aux qubits. Pour ce faire, il est possible d'utiliser le qubit lui-même pour analyser les signaux qu'il reçoit. Cette idée, laquelle est au cœur du domaine de la détection quantique, a notamment été explorée pour améliorer le traitement de l'information quantique dans des qubits supraconducteurs dès 2012 [95]. Depuis, de nombreuses approches ont été développées afin d'étendre le domaine d'application ainsi que de raffiner la précision et l'exactitude d'une telle caractérisation [96, 97, 98, 99, 100]. Par exemple, dans les récents travaux de Hyyppä et co-auteurs [100], un protocole de caractérisation des distorsions micro-ondes possédant une précision inégale a été développé et utilisé afin de démontrer des opérations logiques à un transmon de très haute fidélité (99.98 %) et de très courte durée (7.9 ns).

Pour les signaux à basse fréquence utilisés typiquement dans les opérations à deux qubits [101], il demeure difficile d'éliminer de façon significative l'ensemble des distorsions [94]. Une approche de choix consiste alors à adapter les contrôles utilisés afin de les

rendre moins sensibles aux distorsions que l'on ne peut éliminer. Par exemple, il est possible de largement atténuer l'impact des distorsions à longue durée en appliquant des signaux bipolaires ayant une aire sous la courbe de zéro, alors que les distorsions associées à la seconde moitié des signaux annulent approximativement celles de la première moitié [70, 71].

Somme toute, les exemples de systèmes à deux niveaux et de distorsion des signaux de contrôle nous démontrent l'importance d'obtenir une description précise des dynamiques quantiques du système afin de réaliser des opérations quantiques de haute fidélité. Étant donné qu'effectuer une caractérisation rigoureuse et exacte de tous les degrés de liberté du système quantique est simplement impossible, il devient nécessaire de faire appel à des approches heuristiques. L'idée est alors de trouver une description effective qui est en mesure de reproduire le comportement du système jusqu'à un niveau de précision suffisant. À partir d'un tel modèle et de notre compréhension approximative des dynamiques en jeu, il est ensuite possible d'ajuster avec précision nos contrôles externes afin de réaliser les opérations voulues avec de meilleures fidélités que si l'on considérait un modèle plus simple.

Dans le cadre du projet discuté dans ce chapitre, je propose précisément une telle approche combinant la caractérisation du système quantique à son contrôle. Cette idée a été inspirée par mon projet de recherche à la maîtrise, lequel est présenté dans un article intitulé *Quantum-tailored machine-learning characterization of a superconducting qubit* [102]. L'une des conclusions importantes de ce projet est qu'il est possible d'utiliser une quantité réaliste de données produites par le système étudié afin de construire un modèle décrivant ses dynamiques de manière nettement plus précise que les approches de caractérisation habituelles. Cette approche s'appuie sur les travaux réalisés à l'Université de Berkeley par E. Flurin et ses collaborateurs [103]. Ces travaux démontrent qu'il est possible d'utiliser une approche d'apprentissage automatique initialement développée pour le traitement du langage afin de prédire les dynamiques d'un transmon. En intégrant directement les caractéristiques quantiques du problème dans la conception de l'approche d'apprentissage, nous démontrons qu'il est possible d'améliorer à la fois l'efficacité de l'entraînement et la précision du modèle résultant.

Comme illustré à la Fig. 4.1, l'idée avec ce nouveau projet est d'utiliser ce que nous avons appris en nous attaquant à la tâche de caractérisation précise d'un dispositif quantique afin de compléter une boucle de rétroaction et de directement améliorer les opérations réalisées sur le système quantique. L'utilisation du modèle pour raffiner les opérations quantiques peut prendre diverses formes, par exemple la calibration des paramètres fixes du système (tels que les fréquences d'opération, l'amplitude et la fréquence des impulsions), la conception de nouvelles opérations ou de nouvelles formes de contrôles par ingénierie

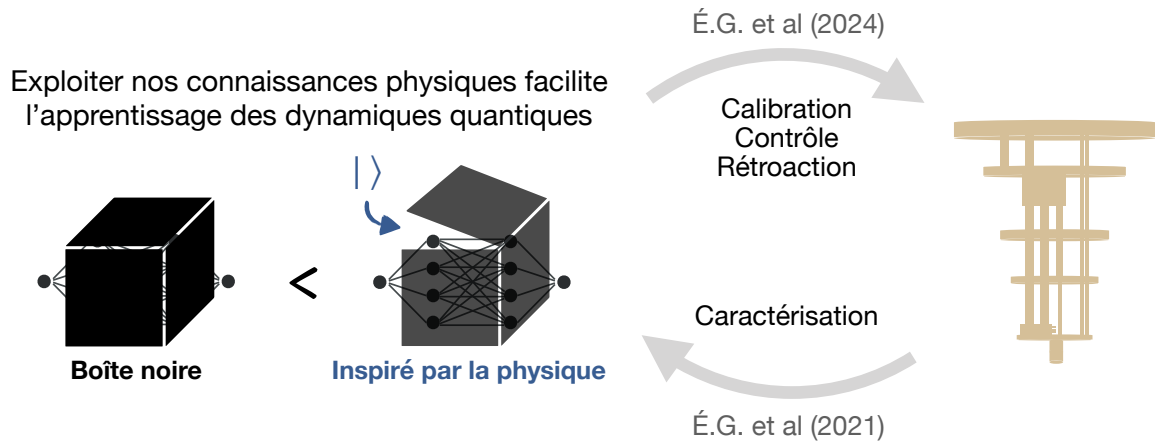


FIGURE 4.1 – Il est possible d'utiliser un réseau de neurones afin d'apprendre précisément la dynamique d'un système quantique à partir de données expérimentales. L'un des résultats clé de mon projet de maîtrise (2021) [102] est que cette caractérisation peut être grandement facilitée et améliorée en adaptant notre approche d'apprentissage aux spécificités du problème physique considéré. Une telle approche inspirée par la physique peut être comparée favorablement à l'approche habituelle traitant le modèle d'apprentissage comme une boîte noire. L'objectif du projet présenté dans ce chapitre (2024) [2] est d'utiliser un tel modèle construit par apprentissage automatique afin de fermer la boucle et d'améliorer les opérations réalisées sur le dispositif quantique.

des impulsions, ou même l'ajustement en temps réel des impulsions avec un contrôle par rétroaction. Ici, je me concentre sur l'utilisation de la théorie du contrôle optimal quantique afin de produire les contrôles réalisant les opérations voulues. Comme nous verrons dans ce qui suit, cette approche est très performante et le fait de la combiner avec un modèle entraîné directement sur des données expérimentales permet d'éliminer son principal inconvénient : le biais de modèle.

4.2 Concepts clés

Dans cette section, j'introduis les deux approches au cœur de ce projet de recherche, soient le contrôle optimal de qubits supraconducteurs et l'apprentissage automatique, en plus de présenter le fort lien unissant ces deux domaines.

4.2.1 Contrôle optimal d'un transmon

Tel qu'illustré schématiquement à la Fig. 4.2(a), la façon principale dont on contrôle les qubits supraconducteurs consiste à générer des signaux micro-ondes à l'aide d'un générateur d'onde arbitraire (AWG), un appareil électronique opérant à température pièce. Les impulsions sont par la suite dirigées dans le réfrigérateur à dilution à l'aide de lignes de transmission, typiquement des câbles coaxiaux, afin d'atteindre les qubits se trouvant à des températures effectives d'environ 10 mK. Tel que décrit à la section 2.4, ces signaux génèrent des opérations à un qubit arbitraires qui sont représentées par des transformations unitaires comme la porte à un qubit $U_{\text{cible}} = \sigma_x$ illustrée à la Fig. 4.2(a).

La théorie du contrôle optimal quantique (QOCT) [104, 105, 6] formalise différentes approches permettant de trouver la forme exacte des contrôles externes à appliquer afin de réaliser une opération quantique voulue avec une fidélité maximale et en un minimum de temps. Cette théorie est fondée sur les outils développés dans le domaine mathématique de la théorie du contrôle [106, 107], lesquels ont dû être adaptés aux particularités de la mécanique quantique [108, 36, 109]. Un domaine actif de recherche du QOCT porte sur la contrôlabilité, à savoir si une opération cible est atteignable en principe étant donné une infidélité et une durée finies [110]. Ici, on s'intéresse plutôt à la conception explicite de ces contrôles, avec pour objectif de trouver la manière optimale de réaliser l'opération en pratique [6].

En raison de la complexité requise pour décrire l'évolution des systèmes quantiques d'intérêt, les méthodes numériques sont au cœur de la conception des contrôles optimaux. La résolution d'un problème de contrôle optimal, lequel est formulé à partir du principe du maximum de Pontryagin (PMP) [106], se résume à trouver la séquence de contrôle qui minimise une fonction de coût étant donné une équation du mouvement et des contraintes fixes, typiquement des conditions frontières. La résolution du PMP est très ardue en pratique étant donné la multitude de contraintes couplées à satisfaire. L'approche par excellence consiste alors à utiliser une approximation au premier ordre du principe du maximum et à résoudre le système d'équations différentielles non linéaires couplées résultant de façon itérative, où la solution testée à une itération donnée est obtenue à partir du gradient de la solution précédente [109]. C'est précisément cette approche d'optimisation par descente du gradient qui est utilisée par la méthode la plus répandue en contrôle optimal quantique : *Gradient-Ascent Pulse Engineering* (GRAPE) [111].

Le fonctionnement de GRAPE et de la grande majorité des approches de contrôle optimal quantique dites à boucle ouverte (c'est-à-dire n'utilisant pas directement le système quantique que l'on veut contrôler) est illustré à la Fig. 4.2(b). Après un paramétrage de la

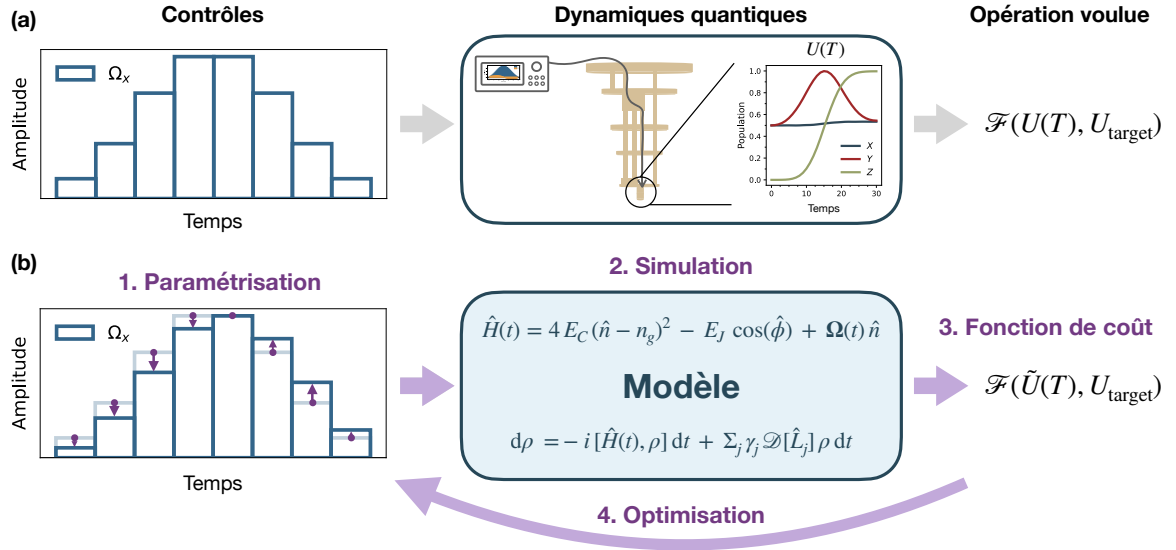


FIGURE 4.2 – Approche standard de contrôle optimal d’un qubit supraconducteur. **(a)** Le contrôle des qubits supraconducteurs passe principalement par la génération d’impulsions micro-ondes à l’aide d’appareils électroniques classiques. Ces signaux se propagent ensuite jusqu’aux qubits afin de générer des dynamiques temporelles non triviales. L’objectif du contrôle quantique consiste alors à concevoir ces impulsions de sorte à produire les opérations quantiques voulues. La qualité de ces opérations peut être quantifiée à l’aide d’une mesure appelée la fidélité de porte moyenne $\mathcal{F}$. **(b)** La conception des impulsions micro-ondes peut être réalisée à l’aide d’une approche de contrôle optimal quantique par boucle ouverte. Cette approche consiste à paramétrer les contrôles d’entrée, à simuler l’évolution du système quantique due à ces impulsions en utilisant un modèle numérique, à calculer une fonction de coût (typiquement la fidélité de porte), puis à utiliser un algorithme d’optimisation afin de minimiser la fonction de coût et ainsi trouver les signaux de contrôle optimaux. Cet algorithme consiste typiquement à réaliser une descente de gradient, comme c’est le cas pour l’approche GRAPE très répandue [111]. De façon cruciale, la performance des contrôles optimaux sur le dispositif quantique réel dépend de la précision avec laquelle le modèle est en mesure de tenir compte des dynamiques du dispositif.

forme des impulsions de contrôle, par exemple un utilisant une amplitude spécifiée à chaque pas de temps, l’idée consiste à utiliser un modèle du système quantique d’intérêt. Ce modèle numérique permet à la fois de simuler l’évolution du système produite par les impulsions données en entrée ainsi que de calculer la dérivée des dynamiques résultantes par rapport aux paramètres de contrôle. La qualité avec laquelle l’impulsion utilisée accomplit l’opération voulue est quantifiée à l’aide d’une fonction de coût, typiquement la fidélité moyenne de l’unitaire résultante. Un algorithme de descente de gradient est enfin utilisé pour mettre à jour la forme de l’impulsion d’entrée de manière à minimiser cette fonction de coût. Cette boucle d’optimisation est répétée jusqu’à la convergence de la solution, typiquement dans un minimum local de l’espace des solutions.

Alors que différentes approches existent afin de raffiner cette optimisation, par exemple

en utilisant des descentes de gradient du deuxième ordre [112, 113] ou en s’assurant de trouver le minimum global de l’espace des solutions [114, 115], l’approche de type GRAPE décrite plus haut fonctionne remarquablement bien en simulation. En effet, il est systématiquement possible de trouver des opérations quantiques avec des précisions limitées seulement par la précision numérique utilisée, et ce, avec des durées de contrôle significativement plus courtes que ce qu’il est possible de faire avec des formes d’impulsions plus simples, lesquelles sont typiquement plates et monochromatiques.

Cependant, cette performance repose de façon cruciale sur le fait que le modèle numérique utilisé décrit l’ensemble des dynamiques du système quantique considéré, et ce, avec une très haute précision. Dans ce contexte, le modèle représente la transformation prenant en entrée les impulsions micro-ondes de contrôle et fournissant en sortie l’état quantique du système produit par ces impulsions. Comme décrit à la section 2.5, ce modèle prend typiquement la forme d’une équation maîtresse de Lindblad où les paramètres de l’Hamiltonien et des taux de perte sont obtenus à l’aide d’expériences de calibration couramment utilisées comme celle de Ramsey. Sans surprise, ce modèle est imparfait et la précision finie de ses prédictions peut significativement limiter la performance des impulsions optimales lorsqu’elles sont appliquées sur le système réel. Il s’avère que cette limitation est prédominante pour les qubits supraconducteurs alors que les réalisations expérimentales du contrôle optimal quantique sont contraintes d’aller au-delà d’une optimisation par boucle ouverte afin de calibrer les impulsions directement sur le système quantique d’intérêt [116, 117, 118].

Afin d’illustrer cet effet, j’effectue l’optimisation d’une rotation de 180° autour de l’axe σ_x , $U_{\text{cible}} = \exp(-i\pi\sigma_x/2)$, sur un modèle de qubit transmon à la Fig. 4.3, puis j’évalue la performance de ce contrôle optimal lorsque la fréquence du qubit est modifiée (via un changement à son énergie Josephson E_J). L’infidélité du contrôle atteint la précision des erreurs numériques lorsque le modèle est parfait dans cette simulation simple où aucun bruit n’est considéré. Cependant, aussitôt qu’un biais de modèle est présent, c’est-à-dire une incompatibilité entre le modèle utilisé lors du contrôle optimal et la dynamique quantique réelle du système (ici la fréquence fondamentale du transmon), l’erreur associée à l’opération quantique croît rapidement, pour atteindre une modeste fidélité de 99.9 % avec un décalage de la fréquence du transmon de moins de 0.5 MHz. Étant donné les fluctuations temporelles importantes de la fréquence des qubits supraconducteurs [119, 120], une telle sensibilité aux paramètres exacts du système s’avère problématique en pratique.

Somme toute, la modélisation physique usuelle des dispositifs quantiques est insuffisante pour réaliser des opérations de très haute fidélité ($> 99.99\%$) avec une approche de contrôle optimal quantique par boucle ouverte. Une solution consiste à raffiner notre modèle des dynamiques du système, une tâche correspondant précisément à la caractérisation

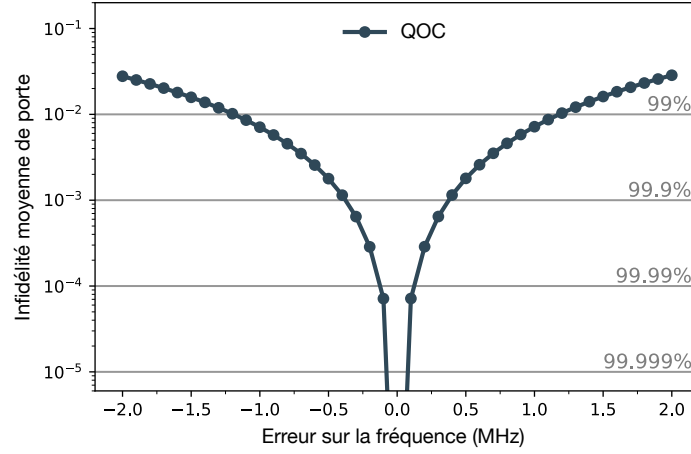


FIGURE 4.3 – Effet du biais de modèle sur une approche de contrôle optimal quantique (QOC) par boucle ouverte. Numériquement, il est possible d’optimiser la forme d’une impulsion micro-onde de 30 ns afin de réaliser une porte logique $\hat{\sigma}_x$ avec une erreur moyenne limitée seulement par la précision numérique utilisée ($< 10^{-12}$). Cependant, en appliquant cette même impulsion sur un modèle simplifié d’un transmon possédant une fréquence d’oscillation légèrement différente de celle supposée lors de l’optimisation, l’infidélité croît rapidement.

quantique discutée dans le chapitre d’introduction. Adoptant une approche variationnelle, il est alors possible de paramétrer la transformation réalisée par notre modèle et d’optimiser ces paramètres afin de décrire les données expérimentales produites par le système de la manière la plus exacte possible. C’est précisément ce que nous cherchons à accomplir dans le cadre de ce projet de recherche à l’aide d’une approche par apprentissage automatique supervisée.

4.2.2 Apprentissage automatique

L’apprentissage automatique est un domaine en pleine effervescence, notamment en raison de l’avenue d’infrastructures de calcul spécialisées supportant l’entraînement de modèles possédant des centaines de milliards de paramètres [121]. Les applications des algorithmes d’apprentissage automatique ne cessent de se multiplier dans divers domaines en raison de leur capacité à résoudre de façon approximative des problèmes exceptionnellement complexes. Le domaine de l’informatique quantique n’y fait pas exception [122, 123, 124, 125, 126]. Alors qu’une introduction rigoureuse au domaine de l’apprentissage automatique dépasse le cadre de cette thèse, j’invite la lectrice et le lecteur intéressés à consulter l’une des excellentes références suivantes [127, 128, 129, 130].

Pour nos besoins, une définition opérationnelle suffisante est que l’apprentissage automatique est une approche permettant d’exploiter les corrélations présentes dans des

ensembles de données afin de réaliser une tâche spécifiée par une fonction de coût. En optimisant les paramètres libres du modèle d'apprentissage afin de minimiser cette fonction de coût, la transformation réalisée par le modèle tend à accomplir la tâche voulue de manière heuristique. La qualité avec laquelle le modèle est en mesure d'accomplir cette tâche dépend à la fois des spécificités du modèle (notamment son expressivité, c'est-à-dire l'espace des transformations qu'il est en mesure d'approximer) et de son entraînement (par exemple la quantité et la qualité des données d'entraînement, l'algorithme d'optimisation utilisé, les hyperparamètres associés à la fonction de coût et à la façon dont les paramètres sont initialisés et modifiés, etc.).

Alors que les architectures de ces modèles d'apprentissage sont constamment en développement, un concept clé pour améliorer la performance du modèle consiste à spécialiser la transformation qu'il réalise aux spécificités du problème considéré. Comme présenté à la section 4.1, c'est ce que nous avons démontré dans mon projet de maîtrise dans le cadre d'une tâche de caractérisation quantique [102]. De façon plus générale, cette idée de spécialisation du modèle d'apprentissage est répandue dans le domaine. On note par exemple le succès des réseaux de neurones convolutifs (CNN) en reconnaissance d'images alors que ceux-ci préservent les corrélations spatiales entre les pixels d'entrée en analysant les données sur une parcelle à la fois [131, 132, 133]. De façon similaire, les réseaux de neurones récurrents (RNN) préservent l'information chronologique des séquences de données en traitant les entrées une à la fois tout en gardant une mémoire des entrées reçues précédemment [134]. Alors que les données ordonnées séquentiellement se retrouvent dans de très nombreuses applications physiques où la dynamique temporelle d'un système est en jeu, c'est sans surprise que les RNNs sont couramment utilisés en informatique quantique [135, 136, 137, 103]. Nous utilisons également ce type d'architecture dans le cadre de mon projet, plus spécifiquement un réseau *Long-Short Term Memory* (LSTM) [138]. Il s'agit du type de RNN le plus répandu alors que sa structure de régulation de l'information permet l'apprentissage de corrélations sur de longues séquences. L'idée au cœur du succès de cette architecture consiste à gérer l'introduction et la suppression des caractéristiques en mémoire à l'aide de structures agissant comme *portes*. Ces transformations font en sorte que l'information soit en mesure de se propager sur de longues séquences tout en évitant que les gradients ne s'estompent trop rapidement pour influencer les sorties futures [127].

Retournant à notre problème de contrôle optimal de qubits supraconducteurs nécessitant une meilleure caractérisation du dispositif quantique, nous sommes en lieu de se demander pourquoi utiliser une approche d'apprentissage automatique alors que l'on dispose d'approches rigoureuses fondées sur des principes physiques pour réaliser cette caractérisation. Les deux raisons principales sont la prévention du problème de biais du modèle ainsi que

la simplification de sa réalisation pratique en réduisant le nombre d'expériences requises.

L'utilisation d'apprentissage automatique permet de grandement atténuer le problème de biais du modèle alors qu'en vertu du théorème d'approximation universelle, les réseaux de neurones sont en mesure d'approximer n'importe quelle fonction continue sur des sous-ensembles euclidiens [139]. Ainsi, la transformation apprise par le modèle d'apprentissage peut être rendue complètement générale et donc agnostique aux approximations physiques typiquement utilisées. En utilisant un modèle d'apprentissage suffisamment expressif pour représenter le comportement complet du dispositif quantique, il est alors possible d'utiliser des données expérimentales produites par ce dispositif afin d'apprendre ses dynamiques sans aucun biais.

La seconde motivation principale pour utiliser l'apprentissage automatique à la place d'une approche fondée uniquement sur la physique du problème est d'obtenir une solution approximative de façon plus efficace, c'est-à-dire nécessitant moins d'expériences en pratique. Comme décrit à la section 4.1, la tâche de caractérisation est très complexe en raison des nombreux degrés de liberté pouvant affecter la dynamique du dispositif quantique que l'on désire contrôler. Alors qu'une caractérisation explicite de chacun de ces degrés de liberté devient rapidement irréalisable en pratique à mesure que la taille du système et/ou la précision requise augmentent, il devient nécessaire d'utiliser une approche heuristique différente. L'approche d'apprentissage automatique est alors particulièrement prometteuse étant donné ses nombreux succès démontrés à résoudre des problèmes de complexité similaire, et ce, sans aucune connaissance des étapes intermédiaires ou sous-composantes du problème à résoudre.

Un exemple important de l'utilité d'une approche heuristique d'apprentissage automatique à la résolution d'un problème de nature quantique se retrouve en tomographie d'états quantiques (QST). Cette tâche consiste à reconstruire l'état quantique d'un système à partir de mesures sur un ensemble de copies de cet état. Afin d'identifier cet état de façon unique, les opérateurs de mesure doivent former une base complète de l'espace d'Hilbert du système, ce qui nécessite un nombre exponentiel d'expériences et de paramètres à estimer, par exemple 4^n pour un système de n qubits. La résolution rigoureuse du problème de QST est simplement insurmontable pour des états à plusieurs corps ($n \gtrsim 10$).

Cependant, il existe des représentations alternatives permettant de décrire de nombreux états quantiques complexes de manière approximative, voire exacte, en profitant de leur structure. L'exemple par excellence est l'utilisation de réseaux de tenseurs [140]. En utilisant un tel réseau comme représentation variationnelle de l'état quantique, il est alors possible d'optimiser ses paramètres afin de reproduire les statistiques de mesures observées de

l'état quantique à reconstruire, et ce, avec un nombre réduit d'opérateurs par rapport à une reconstruction complète [141]. Par exemple, il a récemment été démontré qu'un nombre de copies de l'état quantique augmentant de façon polynomiale était suffisant pour reconstruire, avec une erreur bornée, des états quantiques à plusieurs corps possédant une structure d'intrication locale à l'aide d'une telle approche [142]. Ces mêmes idées s'appliquent à la tomographie de processus quantiques (QPT) où la dynamique temporelle d'une opération est caractérisée [143], une tâche qui s'apparente à celle réalisée dans le cadre de mon projet. Je note cependant que contrairement à la description d'une seule opération en QPT, nous cherchons à caractériser l'ensemble des dynamiques du dispositif pour des contrôles arbitraires. Il est alors nécessaire d'estimer les paramètres dynamiques du système quantique, une tâche souvent décrite comme *l'apprentissage de l'Hamiltonien* [144, 145].

Voyons maintenant les différentes façons d'utiliser l'apprentissage automatique dans le contexte du contrôle optimal de dispositifs quantiques.

4.2.3 Apprentissage automatique pour le contrôle quantique

Au cours des dernières années, l'avenue principalement explorée pour résoudre le problème du biais de modèle en contrôle optimal quantique consiste à concevoir les contrôles externes en interagissant directement avec le système quantique. Cette approche utilise une boucle d'optimisation fermée, laquelle est illustrée à la Fig. 4.4. L'idée de ce contrôle par rétroaction est d'utiliser certaines observations expérimentales afin de mettre à jour la stratégie de contrôle utilisée, puis de tester ces nouveaux contrôles directement sur le système quantique afin d'obtenir de nouvelles observations et d'itérer jusqu'à l'obtention d'une stratégie optimale de contrôle. De nombreuses propositions [146, 147, 148, 149, 150], de même que plusieurs réalisations expérimentales [94, 55, 151], utilisent l'apprentissage par renforcement (RL) [129] pour résoudre ce problème d'optimisation en boucle fermée. Dans ce type d'apprentissage automatique, au lieu d'optimiser les paramètres libres du modèle de sorte à minimiser une fonction de coût évaluée sur un ensemble de données fixe, le modèle procède par essai et erreur en interagissant directement avec le système d'intérêt. Le système est typiquement appelé *environnement* dans ce contexte.

Le principal avantage de cette approche est qu'elle peut être construite de sorte à n'utiliser aucun modèle, c'est-à-dire aucune conception préalable sur la façon dont les actions (contrôles) influencent les observations reçues. Cette approche élimine donc complètement le problème du biais de modèle en n'assumant absolument rien sur la dynamique du système quantique.

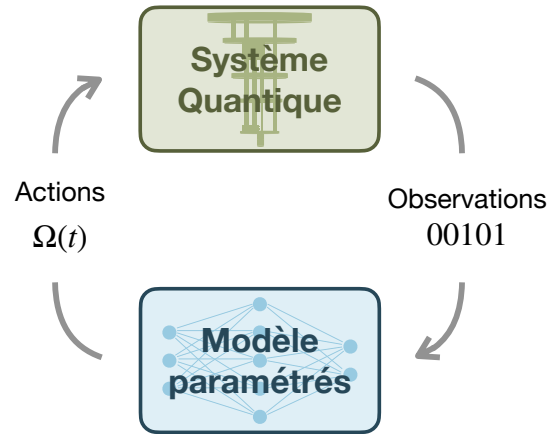


FIGURE 4.4 – Contrôle par rétroaction. Certaines actions avec des particularités bien définies $\Omega(t)$, par exemple des impulsions micro-ondes avec diverses enveloppes temporelles, sont appliquées sur le système quantique d'intérêt. Le système est ensuite mesuré directement afin de produire des observations quantitatives, par exemple des bits de mesures projectives. Ces observations sont utilisées par le contrôleur, lequel prend la forme d'un modèle quelconque avec des paramètres libres tels qu'un réseau de neurones. Le contrôleur modifie alors la stratégie de contrôle utilisée de sorte à minimiser une fonction de coût associée aux observations reçues. L'adaptation de la stratégie de contrôle est typiquement réalisée en combinant une descente de gradients avec un algorithme d'exploration de l'espace de contrôle. L'apprentissage par renforcement (RL) est un outil idéal pour résoudre ce genre de problème de contrôle par rétroaction.

Il est cependant important de noter que le simple fait d'utiliser une approche d'apprentissage par renforcement dite *sans modèle*, telle que l'algorithme appelé *Q-learning* [152] ou l'optimisation de politique proximale (PPO) [153], n'est pas suffisant pour s'affranchir complètement du biais de modèle. En effet, plusieurs approches de RL sans modèle qui ont été proposées se basent en fait sur un modèle physique du système. Par exemple, l'approche présentée dans la Réf. [146] utilise explicitement un Hamiltonien et résout l'équation de Schrödinger afin d'évaluer la qualité des contrôles utilisés. De façon similaire, les auteurs de la Réf. [149] construisent l'état quantique du système à partir de l'information incomplète provenant des observations réalistes, ce qui nécessite un modèle d'équation maîtresse. Une approche pour éviter d'utiliser une modélisation physique est alors de construire l'état quantique du système directement à partir d'observations complètes. Cette approche basée sur la tomographie d'états quantiques a été démontrée dans la Réf. [94]. Elle fonctionne bien pour le contrôle de petits systèmes quantiques, tel qu'un seul qubit, mais est impossible à réaliser en pratique pour contrôler des systèmes quantiques plus complexes étant donné la quantité requise d'expériences.

Je souligne qu'il est possible de rendre l'approche de contrôle par RL réellement exempte de tout biais et applicable à plus grande échelle, comme décrit dans l'excellente Réf. [148].

L'idée consiste à soigneusement concevoir l'expérience réalisée de sorte que les observables quantiques mesurées correspondent directement à la fonction de coût que nous souhaitons minimiser, typiquement l'infidélité d'un état ou d'un processus. Par exemple dans le problème de préparation d'états bosoniques arbitraires, il est possible d'intriquer l'état de la cavité contenant l'état quantique d'intérêt avec un qubit ancillaire de sorte que la mesure du qubit permette d'effectuer la tomographie de Wigner de l'état bosonique [154]. L'obtention de quelques échantillons s'avère suffisante pour estimer la fidélité de l'état bosonique préparé et ainsi accomplir la tâche de contrôle par RL sans biais de modèle [148].

L'utilisation judicieuse de contrôle par RL s'avère très utile afin de calibrer une opération spécifique avec une grande fidélité et les réalisations expérimentales de cette approche avec des qubits supraconducteurs continuent de s'accroître [55, 151, 155]. Cependant, cette approche vient avec un lot d'inconvénients. Tout d'abord, son implémentation nécessite des interactions en temps réel avec le dispositif quantique, ce qui peut s'avérer difficile à réaliser expérimentalement. De plus, l'agent apprenant la stratégie de contrôle optimale agit comme une boîte noire, c'est-à-dire qu'il n'apprend qu'une seule opération quantique et qu'on ne peut typiquement pas utiliser l'information apprise pour effectuer d'autres tâches de contrôle ou de caractérisation. Le dernier inconvénient, mais sans doute le plus important, est que cette approche nécessite un grand nombre d'interactions adaptatives avec le système quantique, sans aucune garantie de converger vers une bonne solution.

Ce besoin en données et en temps expérimental découle du fait que le modèle d'apprentissage apprend à partir de zéro et doit résoudre une tâche très complexe. En effet, le modèle doit à la fois construire une représentation sur la façon dont les actions disponibles influencent la fonction de coût à optimiser, uniquement en se basant sur son expérience des actions et observations reçues jusqu'à ce point, en plus de devoir naviguer l'énorme espace des contrôles afin de trouver la solution optimale. Reconnaisant le fait que ces deux aspects de la tâche d'apprentissage peuvent être complètement séparés, je propose une approche en deux étapes distinctes permettant de simplifier le problème à résoudre. L'objectif est alors d'améliorer l'efficacité avec laquelle nous sommes en mesure de réaliser des opérations quantiques de haute fidélité en pratique.

4.2.4 Contrôle optimal en deux étapes

Comme illustré à la Fig. 4.5, je propose de diviser en deux la transformation passant des observations expérimentales (entrées) à la forme du contrôle optimal (sortie). Tel que détaillé dans l'article scientifique qui suit à la section 4.4, l'approche consiste d'abord à appliquer un ensemble de contrôles différents sur le dispositif quantique d'intérêt et

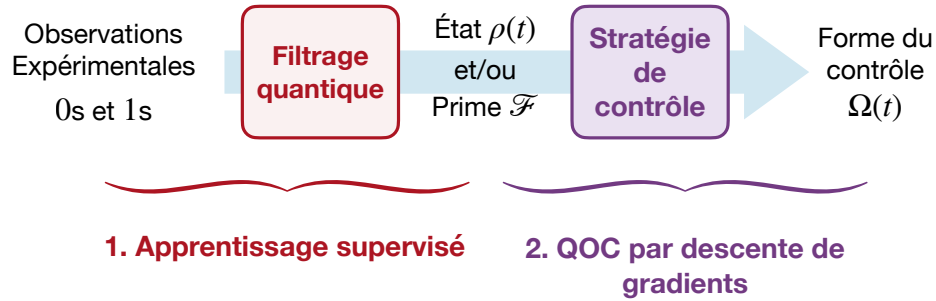


FIGURE 4.5 – Les deux étapes de l’approche proposée permettant de contrôler un dispositif quantique avec une grande fidélité. Plutôt que de chercher à passer directement des observations expérimentales aux contrôles réalisant l’opération voulue de façon optimale (flèche bleue), je propose de séparer cette transformation en deux. Tout d’abord, les dynamiques quantiques du dispositif sont caractérisées à l’aide d’une approche d’apprentissage automatique supervisé (rouge). Par la suite, ce modèle entraîné décrivant le comportement du dispositif est utilisé dans une optimisation par descente de gradients afin de trouver la forme du contrôle réalisant l’opération voulue avec une fidélité optimale (mauve).

d’effectuer des mesures projectives. Cette expérience produit un ensemble de données échantillonnant la transformation subie par le système lors de l’application d’impulsions arbitraires. Alors que chaque entrée de l’ensemble de données (ici des impulsions micro-ondes) est associée à une étiquette (ici le résultat d’une ou de plusieurs mesures projectives), on utilise une approche d’apprentissage automatique dite supervisée afin d’entraîner le modèle paramétré (ici un réseau de neurones récurrent) à prédire les étiquettes avec le plus de précision possible. Cette première étape produit alors un modèle décrivant les dynamiques quantiques du système. En utilisant un modèle d’apprentissage suffisamment expressif et ayant accès à un grand ensemble de données, cette description devrait être plus exacte que la modélisation physique typiquement utilisée. Nous démontrons et quantifions le tout dans l’article qui suit.

La deuxième étape consiste simplement à utiliser le modèle entraîné dans l’une des nombreuses approches de contrôle optimal quantique dont la performance est bien reconnue. Une telle optimisation permet alors d’obtenir la forme des contrôles réalisant les opérations voulues avec une infidélité et une durée minimales tout en respectant les contraintes expérimentales pertinentes. Plus spécifiquement, nous utiliserons une approche par descente de gradients alors qu’il est particulièrement aisé d’obtenir la dérivée première d’une fonction de coût arbitraire par rapport aux paramètres d’un réseau de neurones en profitant de la différentiation automatique [156, 157].

En plus de simplifier la tâche d’apprentissage nécessitant des données expérimentales, cette approche en deux étapes possède plusieurs avantages par rapport aux approches de contrôle par rétroaction. D’abord, le modèle entraîné décrivant les dynamiques du

dispositif lors de l'application de contrôles arbitraires nous permet d'optimiser tout un ensemble d'opérations sans nécessiter de données supplémentaires. En plus d'être en mesure d'optimiser un ensemble universel de portes logiques d'un seul coup, cette caractéristique peut être particulièrement avantageuse afin de réaliser une simulation quantique ou de compiler un algorithme quantique en un circuit de profondeur réduite suite à l'utilisation de portes paramétrées de façon continue [101, 158]. Un autre avantage de l'approche proposée est qu'en plus des applications de contrôle optimal, les données expérimentales de même que le modèle entraîné nous permettent d'obtenir des informations détaillées de caractérisation. Celles-ci peuvent être utiles afin de comprendre et potentiellement d'améliorer le dispositif quantique. Je reviendrai sur ces idées à la section 4.5. Enfin, d'un point de vue pratique, cette approche est très facile à réaliser expérimentalement alors qu'elle ne nécessite aucune rétroaction en temps réel et que la production de l'ensemble de données utilise le dispositif dans son mode d'opération habituel où des impulsions micro-ondes sont envoyées aux qubits et sont suivies par des mesures projectives.

4.3 Résultats principaux

Ce projet de recherche introduit et démontre une approche au contrôle optimal quantique conçue pour pallier le problème de biais de modèle tout en étant aisément réalisable en pratique. À l'aide de simulations numériques réalistes des dynamiques d'un transmon, nous démontrons plus spécifiquement qu'il est possible d'entraîner un réseau de neurones récurrent sur un ensemble de données de taille modeste afin de caractériser le comportement du système quantique avec une excellente précision. De façon plus quantitative, l'utilisation d'un million d'impulsions est suffisante pour entraîner un RNN à prédire l'état quantique d'un qubit supraconducteur avec une erreur moyenne de moins d'un pour cent sur ses valeurs attendues. Un ensemble de données de cette taille ne prend que quelques heures à acquérir sur un dispositif supraconducteur standard, où la grande majorité du temps d'acquisition est dominée par le chargement des impulsions sur les appareils de contrôle.

Par la suite, nous démontrons pour la première fois à notre connaissance l'utilisation d'un réseau de neurones comme modèle dans une boucle ouverte d'optimisation afin de trouver les contrôles optimaux d'un système quantique. Cette approche nous permet de trouver des impulsions micro-ondes réalisant des opérations logiques arbitraires sur le transmon avec des fidélités de plus de 99.98 %. Il s'agit d'une démonstration de principe claire que notre approche combinant l'apprentissage automatique et le contrôle optimal quantique, surnommée *MLQOC*, fonctionne et est prometteuse en pratique.

Afin de motiver l'importance et l'utilité de notre approche, nous comparons également sa performance à l'approche standard de contrôle optimal en boucle ouverte en présence de biais de modèle. Parmi les différentes sources possibles d'un tel biais dans un scénario réaliste de contrôle quantique, nous utilisons la distorsion des impulsions pour illustrer le pouvoir de notre approche. Ce choix nous permet de quantifier aisément l'importance du biais présent dans le système. Nos résultats montrent que contrairement à l'approche standard au contrôle optimal qui est largement affectée par le biais de modèle, la performance de notre approche reste largement inchangée, même en présence d'importantes distorsions. En effet, alors que la fidélité des opérations trouvées par l'approche standard chute rapidement sous le seuil du 99.9 %, la fidélité des opérations du MLQOC demeure à environ 99.98 %. De ce fait, notre approche MLQOC peut mener à un gain considérable en ce qui concerne notre capacité à contrôler un système quantique avec précision dans la situation très répandue où notre modélisation physique du dispositif est simplement insuffisante pour atteindre la fidélité voulue.

Ce gain s'explique par le fait que dans notre approche, le modèle apprend l'ensemble de la transformation expérimentale : des contrôles fournis en entrée (ici la forme des impulsions micro-ondes générées par les appareils électroniques) jusqu'au bit classique de mesure projective qui a été classifié. De ce fait, la transformation apprise inclue les distorsions potentiellement subies par les impulsions, de même que toute autre dynamique présente dans le dispositif et affectant les observables mesurés. Notre modèle d'apprentissage peut donc être considéré comme étant un estimateur non biaisé des dynamiques quantiques et classiques du système étudié. Comme nous le verrons, un tel modèle peut s'avérer très utile dans l'opération et le développement de technologies quantiques.

4.4 Publication [2] : De la caractérisation par apprentissage automatique au contrôle optimal quantique

Fort de l'expérience acquise lors de mon projet de maîtrise, j'ai proposé ce projet de recherche et ai réalisé l'ensemble des travaux numériques et de rédaction menant à l'article scientifique qui suit. Plus spécifiquement, j'ai rédigé l'ensemble du code numérique basé sur la librairie PyTorch [159] permettant l'utilisation de processeurs graphiques. Les principaux algorithmes de la librairie que j'ai construite permettent l'entraînement de réseaux de neurones en utilisant des données basées sur des contrôles micro-ondes et des mesures projectives, l'optimisation de portes logiques quantiques arbitraires par descente de gradients, ainsi que des simulations réalistes des dynamiques de qubits supraconducteurs pilotés

par des impulsions micro-ondes variables dans le temps. J'ai généré toutes les données de simulation, entraîné divers modèles d'apprentissage automatique, optimisé de nombreuses portes logiques, en plus d'analyser les résultats et d'itérer les caractéristiques des approches utilisées afin d'améliorer la performance de notre méthode.

J'ai également poursuivi ma collaboration avec le groupe expérimental d'Irfan Siddiqi à l'Université de Berkeley en Californie dans le cadre de ce projet. Après avoir proposé d'appliquer notre méthodologie sur leur dispositif quantique comportant huit transmons, j'ai travaillé avec Noah Stevenson afin d'adapter l'ensemble des étapes de notre approche à une réalisation expérimentale. Noah a acquis les données sur le qubit supraconducteur étudié et a testé les contrôles optimaux, en plus de contribuer à faciliter la mise en œuvre expérimentale de notre approche dans un contexte plus général. Ces résultats expérimentaux sont toujours en préparation et seront inclus dans une seconde version de l'article qui suit.

Quantum optimal control of superconducting qubits based on machine-learning characterization

Élie Genois,^{1,*} Noah J. Stevenson,² Irfan Siddiqi,^{2,3} and Alexandre Blais^{1,3}

¹*Institut quantique & Département de Physique, Université de Sherbrooke, Québec J1K 2R1, Canada*

²*Quantum Nanoelectronics Laboratory, University of California, Berkeley, Berkeley CA 94720, USA*

³*Canadian Institute for Advanced Research, Toronto, Ontario M5G 1M1, Canada*

(Dated: September 21, 2024)

Implementing fast and high-fidelity quantum operations using open-loop quantum optimal control relies on having an accurate model of the quantum dynamics. Any deviations between this model and the complete dynamics of the device, such as the presence of spurious modes or pulse distortions, can degrade the performance of optimal controls in practice. Here, we propose an experimentally simple approach to realize optimal quantum controls tailored to the device parameters and environment while specifically characterizing this quantum system. Concretely, we use physics-inspired machine learning to infer an accurate model of the dynamics from experimentally available data and then optimize our experimental controls on this trained model. We show the power and feasibility of this approach by optimizing arbitrary single-qubit operations on a superconducting transmon qubit, using detailed numerical simulations. We demonstrate that this framework produces an accurate description of the device dynamics under arbitrary controls, together with the precise pulses achieving arbitrary single-qubit gates with a high fidelity of $\sim 99.99\%$.

I. INTRODUCTION

The precise characterization and control of quantum devices are central to the development of useful quantum technologies. The powerful framework of quantum optimal control (QOC) theory can be used to go from a given description of the quantum dynamics, such as a characterized model, to realizing arbitrary quantum operations with maximal fidelity and minimal duration [1–3]. The successful practical implementation of many QOC approaches, such as GRAPE [4] and Krotov [5], thus rely on having an accurate model of the system dynamics to produce the desired output quantum state or process [6]. In simulation, optimizing the input controls using these methods can routinely yield quantum operations with decoherence-limited or even machine-precision fidelity, and these controls can have much shorter durations compared to what is achievable with simpler, monochromatic and flat, pulse shapes [7–11].

A critical problem with using open-loop QOC approaches in experiments is model bias, since the performance of the resulting controls is intrinsically limited by the underlying model accuracy. Indeed, any mismatch between the model used to describe the quantum evolution and the actual dynamics of the physical system can cause the optimal controls to perform significantly worse in practice. This performance degradation is routinely observed when considering parameter deviations to the assumed model, and has led to the development of robust QOC approaches [12–15]. These approaches sacrifice some control performance for the benefit of robustness under certain model parameter variations. However, beyond being inaccurate with model parameters deviat-

ing from their “true” values due to finite characterization precision and system drifts, the model used for QOC can additionally be incomplete, which also leads to important biases. Out-of-model dynamics typically originate from unaccounted frequency- and power-dependent distortions in the control lines, crosstalk between qubits and control lines, and more generally coupling to spurious modes not included in the system modeling such as material defects, box modes, and neighboring couplers and readout apparatus. Precisely characterizing and modeling the dynamics associated with each of these potential interactions is a challenging task, both experimentally and numerically. There is thus a need for alternative approaches to optimally controlling quantum systems.

The shortcomings of using inaccurate models in the control optimization has led to recent proposals of closed-loop optimization approaches based on reinforcement-learning (RL), which rely on direct interactions between a controller and the quantum system of interest [16–20]. Using such a feedback control approach, the QOC problem can be made model-free, thus alleviating the aforementioned problem of model bias. For example, this approach was applied experimentally to realize high-fidelity single- and two-qubit quantum gates, as well as a quantum error correction stabilization protocol in superconducting quantum devices [21–23]. Although useful for precisely calibrating a specific operation, these model-free approaches possess important drawbacks, notably in terms of sample efficiency, generalizability beyond performing a single quantum operation, ease of experimental implementation given the need for real-time feedback, and performance of the control solutions. For instance, Porotti *et al.* [20] demonstrated that model-based gradient-ascent approaches, such as GRAPE, significantly outperform model-free RL approaches on standard state preparation optimal control tasks, both in terms of state fidelity and data efficiency. This can be

* elie.genois@usherbrooke.ca

understood from the fact that in the model-free setting the task of the learning agent is much more complex, as it needs to both resolve how a given control pulse (action) impacts the state or process fidelity (reward), and learn how to improve that control by navigating the enormous control parameter space. The first part of this task can be recognized as a quantum characterization problem, whereas the second part of finding the optimal control strategy can be directly achieved using one of the many successful QOC approaches, instead of relying on the trial-and-error exploration of the RL approach.

Based on this understanding, we propose in this work to simplify the quantum optimal control problem by breaking it down into two parts that we perform in succession. First, we frame the quantum characterization problem, or model learning problem, as a supervised machine learning (ML) task. We use a parametrized representation to learn a description of the quantum system dynamics directly from experimentally available data. Second, we use that trained model in a gradient-based optimal control loop to find the external controls realizing our target operations. In the following, we demonstrate that this modular and easily implementable approach can yield high-fidelity controls using experimentally realistic data, while providing notable benefits in terms of data efficiency, scalability to complex control problems, and device characterization.

The paper is organized as follows. In Sec. II, we describe our machine-learning-based optimal control approach and its benefits over alternative control approaches. Using detailed numerical simulations, we then present a case study implementation of our approach and demonstrate its performance for realizing high-fidelity arbitrary single-qubit gates in a transmon qubit in Sec. III. In Sec. IV, we analyze the impact of out-of-model dynamics on the optimal control performance, and demonstrate the distinct advantages of our approach in the presence of model bias. We conclude with a short discussion on the relevance of this work in Sec. V.

II. MACHINE-LEARNING BASED QOC

The two main steps of our proposed quantum optimal control approach based on machine-learning characterization are presented schematically in Fig. 1. The characterization task consists of constructing a model that captures the transformation from arbitrary input pulse shapes to the resulting quantum state observables of interest. In this work, we focus on qubit gates and therefore take these observables to be the Pauli matrices on n qubits, which are informationally complete to describe qubit states. We can thus consider that the model we want to construct outputs the full quantum state $\rho(t)$. Note that a more restricted set of observables could be used if acquiring this information is prohibitive, for example in the case of a large Hilbert space, as long as the relevant metric to optimize can be described by the

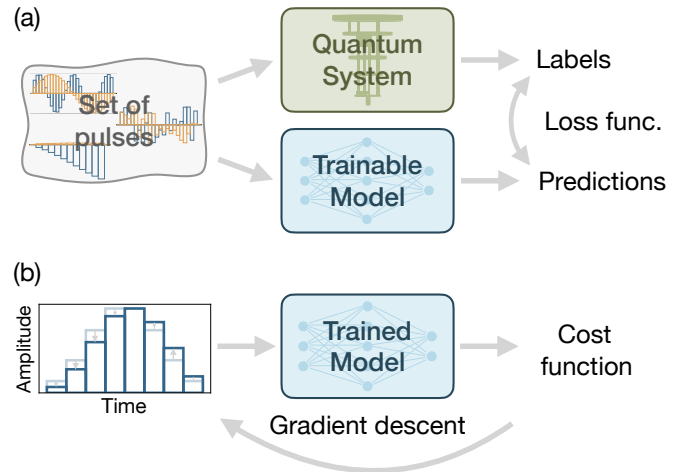


FIG. 1. Optimally controlling a quantum system using machine-learning characterization. **(a)** A supervised machine learning approach is used to characterize the quantum dynamics produced by arbitrary input controls. The dataset is constructed by applying a set of diverse pulses sequentially to the quantum system and obtaining labels of the system observables using (projective) measurements. A parametrized representation (blue box) learns this transformation from input pulse shape to output quantum observables by minimizing a loss function related to the distance between its predictions and the data labels. **(b)** The trained model is used directly in a gradient-based quantum optimal control loop to find the pulses best realizing the quantum operations of interest. The parametrized pulse shape is iteratively updated such as to minimize a cost function associated with a state or process infidelity.

model output observables.

The idea of our approach is to learn this model directly from data taken on the device of interest using a supervised machine learning strategy. To construct the dataset, we simply drive the quantum system sequentially with a large set of different pulses, and perform projective measurements to obtain bits of information about how the quantum state was transformed by each of these pulses. After averaging over a chosen amount of shots, these measurement outcomes form the labels that the model will be trained to predict, using a loss function quantifying the distance between the model predictions and the labels. As detailed in Sec. III, additional physically motivated loss terms can be added to regularize the model's training, as was done in our previous work [24]. We emphasize that this machine-learning training is performed offline such that no real-time feedback between the model and the quantum system is required.

The second step consists of directly using the trained model in a gradient-based optimal control loop to find our control strategy. At this point after the learning stage, the model parameters are fixed and only the parametrized pulse shapes are optimized. Given the auto-differentiable nature of the model and its fast evaluation on graphics processing units (GPUs), these op-

timizations can be performed in less than a minute on standard hardware. Additionally, the optimization can be designed to include arbitrary experimental limitations in the cost function [25], such as to directly find optimal controls that are easily implementable in practice. Importantly, in contrast to direct feedback control approaches, the trained representation is not tied to a specific quantum gate and we can use the trained model to optimize for arbitrary quantum operations which can be described and realized by the same input controls and output observables as the model. For instance, using a model outputting the quantum state of n qubits, we can optimize for arbitrary qubit gates and qubit state preparations. This can be directly generalized to higher dimensional systems such as qudits and bosonic modes, as long as these additional degrees of freedom are observable in the quantum system of interest.

An important advantage of our two-step approach to quantum control is that it is highly modular and customizable. First, having a trained model describing the system dynamics for arbitrary input controls that is fast to evaluate can be very valuable for understanding, calibrating, and improving the quantum system of interest. Indeed, this model can be used as a heuristic digital twin of the device to perform simulations and prototype new protocols. This is especially useful given that the trained model heuristically includes experimentally-relevant imperfections that were learned from data and that are typically not accounted for when modeling the system. Additionally, this model is fully differentiable which means one can replace typically used parameter sweeps or gradient-free approaches with gradient-based optimization of the controls, which are faster and yield better solutions [26, 27].

The second aspect making our approach modular is that the specific architecture used to learn from data, together with the optimal control algorithm used subsequently, can be explicitly engineered to satisfy the specific data, speed and performance requirements of a given quantum system. This is particularly useful to explore the bias-variance tradeoffs of learning an accurate representation of the device dynamics and to perform model selection. For instance, using a black-box learning model to remove any potential bias might require too much data to be trained to reach the desired precision, whereas a fully principled approach based on a physics description might be computationally prohibitive to model or unable to describe the dynamics accurately enough. Using a graybox approach [28] combining physical priors about the device properties together with additional freedom might be optimal in such a realistic scenario, which can be directly realized within our framework. Indeed, our approach allows us to use as much prior knowledge as desired in the learning model, anywhere from a parametrized master equation [24, 29] to a fully general neural network, as long as the model outputs are informationally-complete for quantifying the desired quantum operations fidelity.

We note that approaches of directly combining characterization (system identification) and control have been studied in control theory [30] and applied to quantum systems using a variety of models and data [28, 31–36]. Focused on experimental feasibility and on mitigating the problem of model bias, this work provides a comprehensive framework for applying such an approach to superconducting and other solid-state devices and demonstrates its performance through detailed simulations. In addition, an important distinction of our work is the demonstration for the first time of using a trained neural network as the model for performing open-loop quantum optimal control, which allows us to significantly mitigate any model bias in practice. In contrast to the similar approach by Youssry *et al.* [28], we show that learning an explicit Hamiltonian description of the dynamics is not required for performing arbitrary QOC with high fidelity.

III. TRANSMON SINGLE-QUBIT GATES

We now present a comprehensive case study of our approach applied to realizing fast and high-fidelity microwave single-qubit gates in a fixed-frequency superconducting transmon qubit [37]. We emphasize that our approach is agnostic to the specific physical implementation and could be applied to other platforms following the steps detailed here. The available qubit controls in the chosen superconducting qubit system are microwave pulses generated at room temperature by classical control electronics, which travel down a cryogenic fridge before reaching the qubit. Information about the dynamics resulting from the controls can be acquired via projective measurement of arbitrary qubit operators using standard dispersive qubit readout [38, 39]. Given these input controls and available output observables, the task of interest is to design the external microwave pulses such as to realize arbitrary qubit unitaries with the highest fidelity and minimal duration.

The following four subsections detail the four steps required for achieving this task using our machine-learning based quantum optimal control approach (MLQOC). Namely, the framework consists of (1) acquiring a dataset on the device, (2) training a machine-learning model on this data, (3) performing QOC optimizations using this trained model, and (4) testing the performance of the optimal controls on the device of interest. As a proof of principle of the method, here the data set is generated by numerical simulation of the device (step 1), and the same is true for testing the performance of the optimized controls (step 4). The other two steps remain unchanged when working with an experimental device. In our simulations, we take special care not to make approximations which would artificially simplify the model learning and optimal control tasks considered, in an effort to provide a convincing demonstration of our approach in an experimentally realistic scenario.

A. Simulation dataset

Physical model. We consider a transmon qubit capacitively coupled to a microwave drive line described by the Hamiltonian ($\hbar = 1$) [37]

$$\hat{H}_{\text{tot}}(t) = 4E_C(\hat{n} - n_g)^2 - E_J \cos \hat{\varphi} + \Omega(t)\hat{n} \quad (1)$$

$$= \hat{H}_{\text{trans}} + \Omega(t)\hat{n}, \quad (2)$$

with E_C (E_J) the charging (Josephson) energy, $\Omega(t)$ the applied microwave drive coupling to the transmon charge operator $\hat{n}$, and $\hat{\varphi}$ the canonically-conjugate phase operator. Given that the qubit is operated in the transmon regime, $E_J/E_C = 110$, and that we are interested in quantum operations limited to the qubit subspace, we can safely ignore the gate charge n_g [37]. Throughout this work, we simulate this full Hamiltonian keeping the first 5 eigenstates of $\hat{H}_{\text{trans}}$. We do not perform any rotating-wave approximation (RWA) in order to capture the fast-oscillating dynamics that are relevant for high-fidelity and fast operations [7, 40].

We parametrize the drive pulses $\Omega(t)$ as the input to the arbitrary waveform generator (AWG) that produces in-phase (S_I) and quadrature (S_Q) signals. These signals are then combined to a local oscillator (LO) of frequency ω_{LO} to be up-converted to the desired GHz-frequency signals, a process known as sideband mixing [41, 42]. This pulse parametrization reflects the actual controls of a realistic experiment, allowing our model to capture potential IQ-mixer imperfections and other pulse distortions, in addition to allowing for precise control of the pulses frequency spectrum. In the absence of imperfections in the control electronics and pulse distortion in the control lines, the microwave pulse reaching the qubit is described by

$$\Omega(t) = S_I(t) \cos(\omega_{\text{LO}}t) + S_Q(t) \sin(\omega_{\text{LO}}t). \quad (3)$$

To drive the transmon on resonance at ω_q and avoid IQ-mixer imperfections producing distorted output signals at frequencies close to ω_{LO} , the AWG signals are typically generated by convolving the pulse envelope with an intermediate frequency oscillation at frequency ω_{IF} such that $\omega_{\text{LO}} + \omega_{\text{IF}} = \omega_q$ [42]. We reproduce this experimental condition here using $\omega_{\text{IF}}/2\pi = 100$ MHz and $\omega_{\text{LO}}/2\pi = 6.198$ GHz to reach the qubit frequency at $\omega_q/2\pi \approx 6.298$ GHz.

As shown in Appendix F, when restricting the description to the computational states of the transmon, the effect of the drive $\Omega(t)$ takes the usual form in the rotating frame of the qubit, $\hat{H}_{\text{qubit}} = I(t)\hat{\sigma}_x + Q(t)\hat{\sigma}_y$, for a resonant drive and under the rotating-wave approximation. The real-valued signals I and Q can be expressed as linear combinations of the AWG signals S_I and S_Q . We thus understand how controlling S_I and S_Q allows us to perform arbitrary single-qubit gates. Here, we avoid the approximations mentioned above and simulate the full transmon Hamiltonian Eq. (2) directly using Eq. (3) with arbitrary time-dependent signals S_I and S_Q .

In the AWG as in our simulations, the pulse shapes S_I and S_Q are specified by real amplitudes positioned at a finite number of times during the operation, called *pixels*. We take the size of these pixels to be 1 ns. Importantly, we apply a Gaussian filter to these discrete signals to (i) interpolate between pixels and simulate experimentally-accurate continuous time evolutions beyond a piece-wise constant pulse approximation [7, 27], and (ii) account for the finite bandwidth of the AWG and the filters used in experiments. We use a Gaussian filter standard deviation of 250 MHz.

To simulate the full time dynamics of the system under the application of the drives, we solve the Lindblad master equation (ME)

$$\frac{d\hat{\rho}}{dt} = -i [\hat{H}_{\text{tot}}(t), \hat{\rho}] + \gamma \mathcal{D}[\hat{b}]\hat{\rho} + 2\gamma_{\varphi} \mathcal{D}[\hat{b}^{\dagger}\hat{b}]\hat{\rho}, \quad (4)$$

which accounts for transmon relaxation and dephasing [39]. In this expression, γ (γ_{φ}) is the relaxation (pure dephasing) rate, and $\mathcal{D}[\hat{X}]\hat{\rho} = \hat{X}\hat{\rho}\hat{X}^{\dagger} - \{\hat{X}^{\dagger}\hat{X}, \hat{\rho}\}/2$ the Lindblad dissipator. In the following, we use relatively high relaxation and coherence times of $T_1 = T_2 = 300$ μs , corresponding to $\gamma/2\pi = 3.33$ kHz and $\gamma_{\varphi}/2\pi = 1.67$ kHz. This choice allows us to resolve the finite control precision of our MLQOC approach, beyond the gate fidelity limit set by decoherence. We use the open-source library **dynamiqs** to perform these simulations, which allows us to use GPU acceleration and to efficiently batch the simulations over multiple pulse shapes in parallel [43].

Experiment and dataset generation. As illustrated in Fig. 2(a), the experiment we consider consists of three steps, where (i) the transmon is prepared in one of the six cardinal states of the Bloch sphere, (ii) a microwave drive with a given pulse shape and duration is applied to the qubit, and (iii) the qubit is readout in the basis $\sigma_j \in \{X, Y, Z\}$. This same experiment is repeated N_{shots} times to acquire statistics, where $N_{\text{shots}} \geq 1$, and the label associated with this specific pulse shape and preparation is the qubit expectation value estimate resulting from averaging these shots. We emphasize that we never feed the full quantum state to the model and that this label, i.e. the average measurement result, is a realistic noisy estimate of the true qubit expectation value that the model is trying to predict. State preparation and measurement (SPAM) errors can add noise to these labels. This noise can be made effectively unbiased using standard error mitigation strategies based on inverting the confusion matrix or on measuring half of the shots with the negative measurement operator. For simplicity in evaluating model performances, we did not model the effect of the mitigated SPAM errors which would typically be significantly smaller than shot noise in our simulations ($N_{\text{shots}} = 32$).

To construct the supervised ML datasets, this experiment is repeated for all of the pulse shapes in a chosen set of pulses, while randomizing over the prepared state and measurement axis for each of these pulses. Implementation details for constructing the pulse set such as to

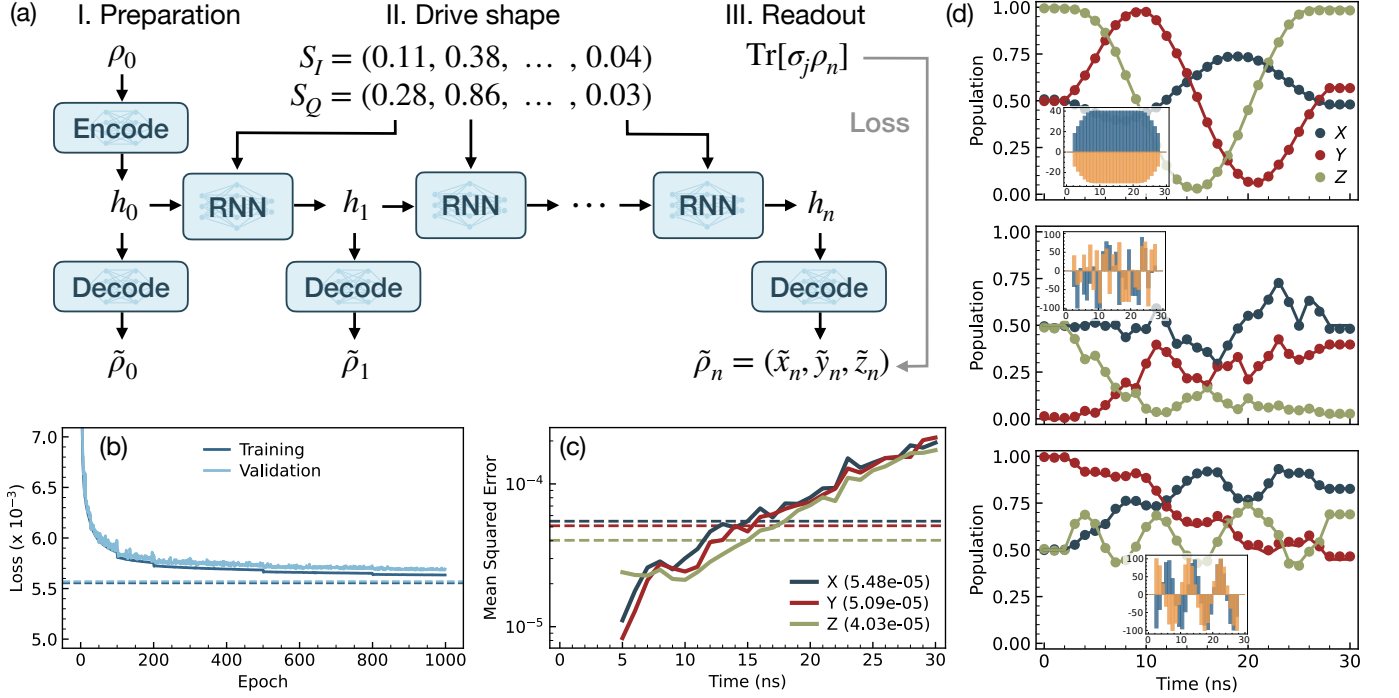


FIG. 2. Learning driven qubit dynamics using supervised machine learning. **(a)** Model architecture. Preserving the structure of the three-step experiment, the model uses a small fully connected neural network (FNN, *Encode*) to map the nominally prepared quantum state (ρ_0) to the initial hidden state representation of the model (h_0). Each pixel of the input pulse shape is fed sequentially to a recurrent neural network cell (*RNN*). A second FNN (*Decode*) is used to produce the predicted quantum state $\tilde{\rho}_n$, here parametrized using Pauli expectation values. During training, the experimental readout outcomes are used as labels and the model parameters are adjusted by minimizing a loss function, here principally composed of a mean squared error (MSE) between the labels and predictions. **(b)** Loss on the training and validation datasets during a supervised LSTM training. All hyperparameters and dataset characteristics are detailed in Appendix B. The dashed lines correspond to the sampling noise floor of the noisy labels for both datasets, see text for details. **(c)** Mean squared error of the predicted quantum state expectation values as a function of pulse duration, averaged over the entire test dataset. Dashed lines correspond to the median over time, with values reported in the legend. The MSE is zero for the first 4 ns because we impose a zero padding of 2 ns at the beginning and end of every pulse. **(d)** Sample model predictions (circles) and true quantum trajectories (full lines) for three input pulse shapes unseen during training. The envelopes of the gaussian flat top (top), random (middle) and sinusoidal (bottom) pulses are shown in inset, with axes corresponding to drive amplitude in MHz and time in ns. Note that during training, the ML model never has access to the true quantum trajectories used here for evaluating model performance in panels (c-d), but only to the noisy labels obtained from $N_{\text{shots}} = 32$ projective measurements.

efficiently sample the control space is presented in Appendix C. We have found that using a combination of random pulse shapes together with physically motivated envelopes, such as Gaussian, sinusoidal, and flat-top envelopes with DRAG components [44] was sufficient for the ML models considered to learn an accurate and generalizable representation of the quantum dynamics.

Putting this data together in the form of supervised learning datasets, each input consists of a one-hot encoded vector $\vec{p}$ describing which of the six cardinal states was prepared, together with a two-dimensional real array $\mathbf{S} = (\vec{S}_I, \vec{S}_Q)$ containing the pulse shape amplitudes at each pixel of the evolution for the two drive quadratures. The output labels are a combination of a one-hot encoded vector $\vec{m}$ capturing which Pauli operator was measured together with a single floating point number representing the associated qubit expectation value estimate $\text{Tr}[\sigma_j \rho(t)]$. Finally, we split this dataset into train-

ing, validation and test sets, as is standard in supervised learning approaches to both avoid over-fitting by selecting the best model from the performance on the validation set, and to obtain unbiased model performance metrics by evaluating the chosen models on a separate test set [45].

B. Model learning

Model architecture. The parametrized representation we use to learn the quantum dynamics from the data described above is the physics-inspired neural-network illustrated in Fig. 2(a). It is principally composed of a recurrent neural network (RNN) architecture [45, 46], which processes the input pulse shapes one pixel at a time such as to preserve the time-ordered structure of the learning problem. The RNN, composed here of a long short-term

memory (LSTM) unit cell [47], is performing the same set of operations at each time step in a recurrent fashion, using the new input together with a hidden state (h) to encode information about the context of that input. This vector allows the model to capture correlations within the input data and to potentially keep a memory of previous inputs. Given our quantum characterization problem where the model needs to output the quantum state at different times, we engineer a direct correspondence between the hidden state of the model h_n the quantum state of the system ρ_n . We thus use an encoding layer, a small fully-connected neural network, to map the initially prepared quantum state to the initial hidden state of the model, and use a similar decoding layer to map this hidden state back to the qubit state prediction $\hat{\rho}(t)$.

To have a model that can efficiently train on a finite and realistic dataset, we have designed the model architecture to explicitly preserve most of the structure of the physical problem at hand. Nonetheless, the model is general enough to go significantly beyond the usual assumptions of a Markovian, single-mode description of the qubit dynamics. As demonstrated in Sec. IV, this choice is motivated by our objective of obtaining high-fidelity controls in scenarios where the physical description used in typical open-loop control approaches is insufficient. We emphasize that many other architecture choices could be directly used within our framework, such as a transformer model [48] or a parametrized master equation.

To train the model, we use a loss function principally composed of a Mean Squared Error (MSE) loss between the model predictions and the labels, as is common for regression tasks in supervised machine learning. Additionally, following Ref. [24], we expand the loss function with terms assuring that the RNN outputs are valid quantum states, i.e. that they are positive, and that the initial state predicted by the model corresponds to the known prepared state of the qubit. These physically motivated loss terms help regularize the model training, as it explores the parameter space of valid representations more efficiently [49]. An explicit expression for the full loss function, together with the hyperparameters used, can be found in Appendix B.

Quantum characterization results. A typical training of the RNN model is presented in Fig. 2(b) on a training dataset comprising of 3.2 million pulse shapes of maximal duration of 30 ns, each measured for 32 shots. As demonstrated in Appendix A, similar performances can be reached using significantly less data. Both the training and validation MSE losses reach a value close to the sampling noise floor of about 5.6×10^{-3} , which is obtained by computing the same MSE metric assuming perfect knowledge of the qubit expectation values, as given by the simulation data. This imprecision is significant as it corresponds to an average distance between the noisy labels and the model output qubit state probabilities of about 7%. Such a high noise floor raises the question of whether the ML model is learning an accurate description of the quantum dynamics.

Leveraging the fact we are working with simulation data, we can go beyond the experimentally-available noisy labels and compare the model predictions directly to the ground truth quantum trajectories obtained from simulation. In Fig. 2(d), we can qualitatively observe the high accuracy of the trained model at describing the qubit dynamics for various input pulse envelopes, as the model predictions (dots) match closely the true trajectories (full lines). This agreement is quantified in Fig. 2(c), showing a median MSE of about 5×10^{-5} across the three measurement axes and over different pulse lengths on the previously unseen test dataset. This order of magnitude improved accuracy on the predicted outcome probabilities indicate that, remarkably, the ML model is able to use the noisy labels that one can obtain experimentally in order to learn the mapping from arbitrary pulse shapes to highly accurate quantum dynamics. The error associated with the model predictions is lowest at the initial time where we have the most accurate information about the qubit state given the finite number of prepared states shared amongst all experiments. The finite model precision for evolving the quantum state for every pixel then leads to a linear increase in the prediction error over time, which is translated into a quadratic increase on the MSE in Fig. 2(c). This behavior is expected for any model with a finite accuracy, and is akin to evolving the state using the ME with slightly inaccurate model parameters, or to numerically solving a differential equation with a finite precision solver.

The accuracy of the ML model can be significantly improved by extending the training dataset, for example using more shots such as to construct more precise labels, and performing more experiments with different pulse shapes such as to explore more of the control space. As we will show next, the amount and quality of the training data used here is sufficient to obtain high-fidelity controls from the model. Additional results exploring how the quality and quantity of training data impact the learned representation and the performance of the resulting optimal controls are presented in Appendix A.

The neural network model takes about an hour to train on a standard Nvidia RTX 3080 GPU. This time could be significantly reduced if needed for a specific implementation, for example by pre-training the model on simulation data before training on the experimental measurements. This timescale is comparable to the experimental data acquisition time for building the datasets using transmon qubits, which is a few hours per million pulse shapes with 32 shots, including the overhead of pulse sequence uploading on the control electronics.

C. Quantum Optimal Control

Having demonstrated that a RNN can be used to accurately learn quantum dynamics from experimentally realistic data, we now present how one can successfully use that representation to perform quantum optimal control.

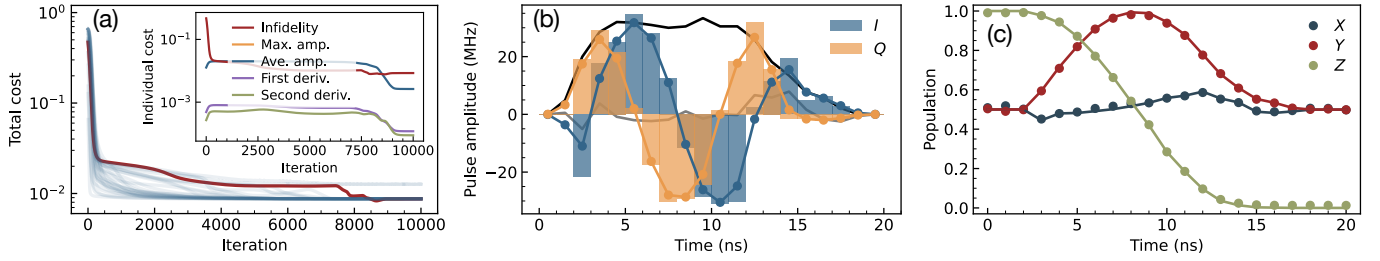


FIG. 3. Quantum optimal control of a 20 ns $R_X(\pi)$ gate on a trained RNN. The optimal MLQOC pulse fidelity is 99.994 % on the 5-level transmon. (a) Parallel optimization of 30 randomly initialized pulse shapes. The optimization of the best performing pulse is highlighted in red, with the inset showing its decomposition into individual costs. (b) Resulting raw (bars) and gaussian filtered (dots with line) optimal pulse shape for the two drive quadratures I and Q . (c) Predicted (dots) and true (lines) quantum state evolution when applying the optimal pulse shape on the initial state $|1\rangle$.

We focus on single-qubit quantum gates and optimize the external controls such as to maximize the average gate fidelity of a given unitary U_{target} . As schematically illustrated in Fig. 1(b), the optimization procedure simply consists of making each pixel of the input pulse a free parameter and computing the evolution of the quantum system due to these controls using the trained ML model with fixed internal parameters. By evolving a set of initial quantum states forming a unitary 2-design, one can compute the average gate fidelity of the arbitrary black-box quantum channel [50, 51], which is represented here by a neural network. For the case of one qubit, the six cardinal states used as preparations in our experiment form a unitary 2-design [52]. Finally, an optimization algorithm is used to update the input controls such as to maximize the fidelity of the operation of interest.

Making use of the fact that the RNN is fully auto-differentiable, we employ a gradient-based approach and design the cost function to include a multitude of experimentally-relevant constraints, without having to analytically derive an expression for their gradient. In addition to the cost associated with the average gate infidelity of the operation of interest, we use regularizing costs that yield smooth and low amplitude controls, which are desirable in the experiment to mitigate crosstalk and pulse distortion effects [25, 53]. In particular, we use a cost term proportional to the average absolute pulse amplitude to penalize for unnecessarily strong pulses, a cost for amplitudes $|\Omega|/2\pi > 100$ MHz, together with terms proportional to the first and second derivatives of the pulse shapes to find smooth controls with a limited frequency spectrum. The full cost function and optimization parameters used here are presented in Appendix D.

A typical control optimization on the trained RNN is shown in Fig. 3(a) for the case of a 20 ns $R_X(\pi)$ transmon qubit gate. Benefiting from the fast forward and backward evaluation of the RNN, 30 randomly initialized pulses are optimized in parallel on a single GPU card in under a minute. This batched optimization and ability to perform a large number of iterations significantly mitigates the convergence issues that many GRAPE-like ap-

proaches face. As a result, we can use the trained model to obtain optimal controls for a wide variety of gates on a very short timescale, which can be used to create a universal gate set of optimal pulses or to compile continuously parametrized gates.

The optimal control found by our MLQOC approach yields a 99.994 % average gate fidelity on the full 5-level transmon model. The pulse shape and resulting qubit dynamics are illustrated in Fig. 3(b-c). Importantly, this optimal pulse is smooth and easily implementable in practice with limited bandwidth electronics. The main sinusoidal oscillations seen in the optimal pulse correspond to the expected intermediate frequency of about $2\pi \times 100$ MHz. Deconvolving this intermediate frequency, we obtain a smooth envelope in the quadrature effectively driving σ_X (black line), together with a small orthogonal component (gray line) necessary to adjust the Z phase of the resulting unitary and to mitigate leakage to the $|2\rangle$ state [44]. As shown in panel (c), the pulse produces coherent dynamics on the qubit that resemble the one of standard sinusoidal and Gaussian control pulses, both according to the RNN model (dots) and the true master equation dynamics (lines).

D. MLQOC performance

An important advantage of our approach is that the trained model is not tied to a specific set of gates and can be used to optimize for arbitrary quantum operations captured by the model inputs and outputs, here any single-qubit unitary. As a demonstration, we use the same trained RNN model presented in Fig. 2 to optimize for a variety of gates and compile the results in Table I. We first optimize for π and $\pm\pi/2$ rotations around the three axes of the Bloch sphere. This gate set, together with the identity operation I realized by applying no drive, generates the full single-qubit Clifford group. We obtain an average fidelity of about 99.99 % for this gate set, demonstrating that finite-precision model trained on realistic data can yield high-fidelity quantum operations. We have also optimized 100 unitaries sam-

SQ gate	Duration (ns)	Fidelity (%)
$R_{X,Y,Z}(\pi)$	20	99.993
$R_{X,Y,Z}(\pm\pi/2)$	15	99.983
Haar random	20	99.984

TABLE I. Gate performance of the MLQOC approach, evaluated on the full 5-level transmon. The trained model can be used to find arbitrary Clifford and Haar random unitaries with high fidelity. Reported fidelities are the average over different axes (X , Y , Z) for Clifford gates and the median of 100 uniformly sampled gates for the Haar random gates. The transmon coherence limits are 99.9975% and 99.9967% for 15 ns and 20 ns gates, respectively. Note that the gate durations include a 4 ns zero padding to avoid gate bleed-through.

pled uniformly at random from the Haar measure [4] and obtained gates with similar performances. This result further demonstrates that the trained RNN representation is accurately capturing arbitrary dynamics of the qubit subspace of our transmon model such that it can be used directly for performing QOC.

The fact that the trained ML model does not fully reach the coherence-limited average gate fidelities of $\sim 99.997\%$ can be attributed to the finite precision of the heuristic that the RNN learns to represent the quantum dynamics. As shown in Appendix A, reducing shot noise and increasing the training dataset size does not significantly improve the control fidelity results, and important stochasticity in the performance of similar models remains, even for trained models performing the quantum characterization task significantly better than the one presented in Fig. 2. Given that it is possible to train a model that would yield coherence-limited gates, such as a model approaching the master equation used to generate the data, it would be interesting to try refining the training data by sampling around the optimal pulses or to explore the performance of different machine-learning architectures. Despite this limitation, we show in the next section that our MLQOC approach can significantly outperform open-loop QOC approaches in practical scenarios by alleviating the problem of model bias.

IV. COMPARISON WITH OPEN-LOOP QOC

Having presented proof-of-principle results of our approach on transmon single-qubit gates, we now demonstrate the advantage of using MLQOC in realistic experimental settings where model bias will necessarily be present. In particular, we show that the MLQOC approach studied here achieves significantly better quantum gates than typical open-loop QOC methods under realistic model bias. This improvement is achieved by learning an accurate representation of the device dynamics from data, and using that model to design high-fidelity controls.

We consider model bias solely coming from out-of-

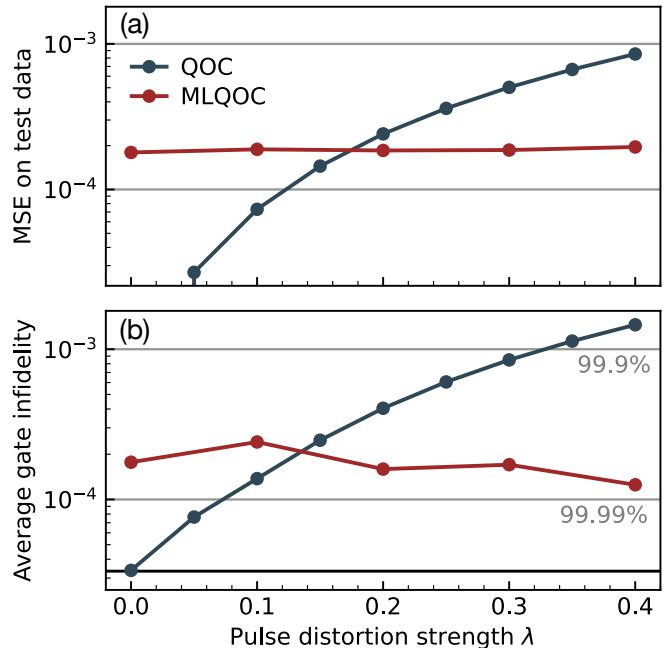


FIG. 4. Demonstrating the usefulness of the proposed MLQOC approach under a realistic source of model bias in the form of random pulse distortions. (a) Predicting the quantum state dynamics of the test dataset using the master equation model (blue) and RNN models trained on experimentally-available data containing 1 million pulse shapes, each sampled 32 shots (red). (b) Performance of the 20 ns $R_X(\pi)$ gates found by optimizing on the models of panel (a), as evaluated on the full transmon model. The gate coherence limit is indicated by a black line.

model dynamics, i.e. effects that are present in the quantum device but not in the physical modeling used by the open-loop QOC approach. We take these dynamics to be unaccounted pulse distortions in the control lines, which produce discrepancies between the pulse shape the qubit receives and the one we model. Note that considering pulse distortions is simply a choice for demonstrating the performance of MLQOC in the presence of model bias, although this choice is motivated by realistic experimental scenarios with superconducting qubits [55–58]. The methodology used to generate the random pulse distortions is detailed in Appendix E.

In Fig. 4, we present the performance of the open-loop (blue) and machine-learning-based (red) QOC approaches at performing the two tasks identified within our framework, namely the quantum characterization as reported by the MSE in panel (a) and the optimal control as reported by the obtained average gate infidelity in panel (b), both as a function of the amplitude of the random pulse distortions λ . Scaling this amplitude λ directly translates into scaling the magnitude of the model bias. For instance, a distortion of $\lambda = 0.2$ corresponds to pulses deviating on average by only 3% from their original shape, which is less than 0.3 MHz.

Since the model used for the open-loop QOC approach

is the full master equation used to describe the transmon qubit, it perfectly describes the quantum system when no model bias is present ($\lambda = 0$), and achieves coherence-limited single-qubit gates. However, as we scale the amplitude of the pulse distortions ($\lambda > 0$), the mean-squared error for predicting the true quantum states of our test dataset grows sharply. Consequently, the controls optimized using the ME model yields single-qubit gate fidelities that rapidly drop below 99.9 % when amplifying the out-of-model dynamics, see panel (b).

In contrast, our proposed ML framework is designed to learn the entire transformation describing our control of the quantum system, from the input pulses of the AWG all the way to the measurement bits we extract, which includes pulse distortions as well as any other dynamics of the quantum device that can affect our measurements. As shown in panel (a), the ML model precision is then virtually unaffected by the added pulse distortions, simply because these dynamics are present in the data it learns from. Using these trained models in panel (b) then allows us to find optimal controls that remain close to 99.99 % fidelity, even under important pulse distortions.

These results constitute a clear demonstration that in a realistic scenario where model bias limits our ability to perform open-loop QOC, our approach of employing ML can lead to a considerable gain in quantum control fidelity. We note that even in a well characterized system where our physical model yields similar quantum characterization performance to the trained ML model, our MLQOC approach can still offer an advantage. This is the case here for $\lambda = 0.2$ where the MSE are similar for both approaches, see panel (a), but the resulting fidelity is superior for MLQOC, see panel (b). This is explained by the fact that the trained model can capture the full device dynamics and thus provide optimal pulses that account for spurious effects such as pulse distortions. As opposed to the physical ME model, the ML model is providing an effectively unbiased estimator of the true quantum dynamics, which for the same level of precision can yield better performing controls in practice.

V. CONCLUSION

In this work, we have presented an experimentally simple two-step approach to quantum control aimed at successfully realizing optimal quantum operations in practice. We tackled the problem of model bias by learning a parametrized representation of the system dynamics directly from experimentally available data, before using that representation to find optimal controls. Using accurate transmon simulations, we have demonstrated that our machine-learning based quantum optimal control (MLQOC) approach can leverage a practical amount of experimental data to provide a precise characterization of the qubit dynamics and produce high-fidelity ($\sim 99.99\%$) controls for arbitrary single-qubit gates, thus going significantly beyond realizing a single operation with high fidelity.

The MLQOC framework promises to enable many quantum optimal control applications where model bias is an important bottleneck, without necessitating any real-time feedback control or having to rely on data-intensive blackbox optimizations of single operations. Beyond experimental ease of implementation, our approach is applicable to a wide range of quantum control problems because its modular features can be adapted to the specificities and requirements of a given quantum system. Importantly, our method of reducing the complexity of the learning task to only the quantum characterization, together with the use of powerful gradient-based QOC techniques, suggests it might be more data efficient and more scalable to complex control problems than model-free approaches. It will be interesting to explore how these approaches perform as we expand towards optimally controlling multi-qubit entangling and parallel operations.

ACKNOWLEDGMENTS

We thank Noah Goss, Pierre Guilmin, Ronan Gauthier, and Boris Varbanov for insightful discussions. This work was undertaken thanks in part to funding from NSERC, the Canada First Research Excellence Fund and the Ministère de l'Économie et de l'Innovation du Québec.

-
- [1] A. P. Peirce, M. A. Dahleh, and H. Rabitz, Optimal control of quantum-mechanical systems: Existence, numerical approximation, and applications, *Phys. Rev. A* **37**, 4950 (1988).
 - [2] J. Werschnik and E. Gross, Quantum optimal control theory, *Journal of Physics B: Atomic, Molecular and Optical Physics* **40**, R175 (2007).
 - [3] C. P. Koch, U. Boscain, T. Calarco, G. Dirr, S. Filipp, S. J. Glaser, R. Kosloff, S. Montangero, T. Schulte-Herbrüggen, D. Sugny, and F. K. Wilhelm, Quantum optimal control in quantum technologies. strategic report on current status, visions and goals for research in europe, *EPJ Quantum Technology* **9**, 10.1140/epjqt/s40507-022-00138-x (2022).
 - [4] N. Khaneja, T. Reiss, C. Kehlet, T. Schulte-Herbrüggen, and S. J. Glaser, Optimal control of coupled spin dynamics: design of NMR pulse sequences by gradient ascent algorithms, *Journal of Magnetic Resonance* **172**, 296 (2005).
 - [5] V. Krotov, *Global Methods in Optimal Control Theory*

- (CRC Press, 1995).
- [6] U. Boscain, M. Sigalotti, and D. Sugny, Introduction to the pontryagin maximum principle for quantum optimal control, *PRX Quantum* **2**, 030203 (2021).
 - [7] F. Motzoi, J. M. Gambetta, S. T. Merkel, and F. K. Wilhelm, Optimal control methods for rapidly time-varying hamiltonians, *Phys. Rev. A* **84**, 022307 (2011).
 - [8] D. J. Egger and F. K. Wilhelm, Optimized controlled-z gates for two superconducting qubits coupled through a resonator, *Superconductor Science and Technology* **27**, 014001 (2013).
 - [9] K. Rojan, D. M. Reich, I. Dotsenko, J.-M. Raimond, C. P. Koch, and G. Morigi, Arbitrary-quantum-state preparation of a harmonic oscillator via optimal control, *Phys. Rev. A* **90**, 023824 (2014).
 - [10] M. H. Goerz, F. Motzoi, K. B. Whaley, and C. P. Koch, Charting the circuit QED design landscape using optimal control theory, *npj Quantum Information* **3**, 1 (2017).
 - [11] L. S. Theis, F. Motzoi, and F. K. Wilhelm, Simultaneous gates in frequency-crowded multilevel systems using fast, robust, analytic control shapes, *Phys. Rev. A* **93**, 012324 (2016).
 - [12] J. H. Wesenberg, Designing robust gate implementations for quantum-information processing, *Phys. Rev. A* **69**, 042323 (2004).
 - [13] X. Wang, L. S. Bishop, J. Kestner, E. Barnes, K. Sun, and S. Das Sarma, Composite pulses for robust universal control of singlet-triplet qubits, *Nature Communications* **3**, 10.1038/ncomms2003 (2012).
 - [14] D. Buterakos, S. Das Sarma, and E. Barnes, Geometrical formalism for dynamically corrected gates in multiqubit systems, *PRX Quantum* **2**, 010341 (2021).
 - [15] A. Koswara, V. Bhutoria, and R. Chakrabarti, Quantum robust control theory for hamiltonian and control field uncertainty, *New Journal of Physics* **23**, 063046 (2021).
 - [16] M. Bukov, A. G.-R. Day, D. Sels, P. Weinberg, A. Polkovnikov, and P. Mehta, Reinforcement Learning in Different Phases of Quantum Control, *Physical Review X* **8**, 031086 (2018).
 - [17] S. Borah, B. Sarma, M. Kewming, G. J. Milburn, and J. Twamley, Measurement-based feedback quantum control with deep reinforcement learning for a double-well nonlinear potential, *Phys. Rev. Lett.* **127**, 190403 (2021).
 - [18] V. V. Sivak, A. Eickbusch, H. Liu, B. Royer, I. Tsioutsios, and M. H. Devoret, Model-free quantum control with reinforcement learning, *Phys. Rev. X* **12**, 011059 (2022).
 - [19] R. Porotti, A. Essig, B. Huard, and F. Marquardt, Deep reinforcement learning for quantum state preparation with weak nonlinear measurements, *Quantum* **6**, 747 (2022).
 - [20] R. Porotti, V. Peano, and F. Marquardt, Gradient-ascent pulse engineering with feedback, *PRX Quantum* **4**, 030305 (2023).
 - [21] Y. Baum, M. Amico, S. Howell, M. Hush, M. Liuzzi, P. Mundada, T. Merkh, A. R. Carvalho, and M. J. Biercuk, Experimental deep reinforcement learning for error-robust gate-set design on a superconducting quantum computer, *PRX Quantum* **2**, 040324 (2021).
 - [22] V. V. Sivak, A. Eickbusch, B. Royer, S. Singh, I. Tsioutsios, S. Ganjam, A. Miano, B. L. Brock, A. Z. Ding, L. Frunzio, S. M. Girvin, R. J. Schoelkopf, and M. H. Devoret, Real-time quantum error correction beyond break-even, *Nature* **616**, 50–55 (2023).
 - [23] L. Ding, M. Hays, Y. Sung, B. Kannan, J. An, A. Di Paolo, A. H. Karamlou, T. M. Hazard, K. Azar, D. K. Kim, B. M. Niedzielski, A. Melville, M. E. Schwartz, J. L. Yoder, T. P. Orlando, S. Gustavsson, J. A. Grover, K. Serniak, and W. D. Oliver, High-fidelity, frequency-flexible two-qubit fluxonium gates with a transmon coupler, *Phys. Rev. X* **13**, 031035 (2023).
 - [24] É. Genois, J. A. Gross, A. Di Paolo, N. J. Stevenson, G. Koolstra, A. Hashim, I. Siddiqi, and A. Blais, Quantum-tailored machine-learning characterization of a superconducting qubit, *PRX Quantum* **2**, 040355 (2021).
 - [25] N. Leung, M. Abdelhafez, J. Koch, and D. Schuster, Speedup for quantum optimal control from automatic differentiation based on graphics processing units, *Phys. Rev. A* **95**, 042318 (2017).
 - [26] J. Nocedal and S. Wright, *Numerical Optimization*, Springer Series in Operations Research and Financial Engineering (Springer New York, 2006).
 - [27] S. Machnes, E. Assémat, D. J. Tannor, and F. K. Wilhelm, Gradient optimization of analytic controls: the route to high accuracy quantum optimal control, *Physical Review Letters* **120**, 150401 (2018).
 - [28] A. Youssry, Y. Yang, R. J. Chapman, B. Haylock, F. Lenzi, M. Lobino, and A. Peruzzo, Experimental graybox quantum system identification and control, *npj Quantum Information* **10**, 1 (2024).
 - [29] S. Krastanov, K. Head-Marsden, S. Zhou, S. T. Flammia, L. Jiang, and P. Narang, Unboxing quantum black box models: Learning non-markovian dynamics (2020), arXiv:2009.03902.
 - [30] P. Albertos and A. Sala, eds., *Iterative Identification and Control* (Springer, London, 2002).
 - [31] R.-B. Wu, B. Chu, D. H. Owens, and H. Rabitz, Data-driven gradient algorithm for high-precision quantum control, *Phys. Rev. A* **97**, 042122 (2018).
 - [32] S. Krastanov, S. Zhou, S. T. Flammia, and L. Jiang, Stochastic estimation of dynamical variables, *Quantum Science and Technology* **4**, 035003 (2019).
 - [33] Q.-M. Chen, X. Yang, C. Arenz, R.-B. Wu, X. Peng, I. Pelczer, and H. Rabitz, Combining the synergistic control capabilities of modeling and experiments: Illustration of finding a minimum-time quantum objective, *Phys. Rev. A* **101**, 032313 (2020).
 - [34] N. Wittler, F. Roy, K. Pack, M. Werninghaus, A. S. Roy, D. J. Egger, S. Filipp, F. K. Wilhelm, and S. Machnes, Integrated tool set for control, calibration, and characterization of quantum devices applied to superconducting qubits, *Phys. Rev. Appl.* **15**, 034080 (2021).
 - [35] H. Ball, M. J. Biercuk, A. R. R. Carvalho, J. Chen, M. Hush, L. A. D. Castro, L. Li, P. J. Liebermann, H. J. Slatyer, C. Edmunds, V. Frey, C. Hempel, and A. Milne, Software tools for quantum control: improving quantum computer performance through noise and error suppression, *Quantum Science and Technology* **6**, 044011 (2021).
 - [36] A. Fyrrillas, O. Faure, N. Maring, J. Senellart, and N. Belabas, Scalable machine learning-assisted clear-box characterization for optimally controlled photonic circuits (2024), arXiv:2310.15349 [physics.optics].
 - [37] J. Koch, T. M. Yu, J. Gambetta, A. A. Houck, D. I. Schuster, J. Majer, A. Blais, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf, Charge-insensitive qubit design derived from the cooper pair box, *Phys. Rev. A* **76**, 042319 (2007).

- [38] A. Blais, R.-S. Huang, A. Wallraff, S. M. Girvin, and R. J. Schoelkopf, Cavity quantum electrodynamics for superconducting electrical circuits: an architecture for quantum computation, *Physical Review A* **69**, 062320 (2004).
- [39] A. Blais, A. L. Grimsmo, S. M. Girvin, and A. Wallraff, Circuit quantum electrodynamics, *Rev. Mod. Phys.* **93**, 025005 (2021).
- [40] R. Shillito, J. A. Gross, A. Di Paolo, É. Genois, and A. Blais, Fast and differentiable simulation of driven quantum systems, *Physical Review Research* **3**, 033266 (2021).
- [41] P. Krantz, M. Kjaergaard, F. Yan, T. P. Orlando, S. Gustavsson, and W. D. Oliver, A Quantum Engineer's Guide to Superconducting Qubits, *Applied Physics Reviews* **6**, 021318 (2019).
- [42] M. Baur, *Realizing quantum gates and algorithms with three superconducting qubits*, Ph.D. thesis, ETH Zurich (2012).
- [43] P. Guilmin, R. Gautier, A. Bocquet, and É. Genois, dynamiqs: an open-source python library for gpu-accelerated and differentiable simulation of quantum systems (2024), dynamiqs.org.
- [44] F. Motzoi, J. M. Gambetta, P. Rebentrost, and F. K. Wilhelm, Simple Pulses for Elimination of Leakage in Weakly Nonlinear Qubits, *Physical Review Letters* **103**, 110501 (2009).
- [45] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning* (MIT Press, 2016) <http://www.deeplearningbook.org>.
- [46] Z. C. Lipton, J. Berkowitz, and C. Elkan, A critical review of recurrent neural networks for sequence learning (2015), [arXiv:1506.00019](https://arxiv.org/abs/1506.00019) [cs.LG].
- [47] S. Hochreiter and J. Schmidhuber, Long Short-Term Memory, *Neural Computation* **9**, 1735 (1997).
- [48] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. u. Kaiser, and I. Polosukhin, Attention is All you Need, in *Advances in Neural Information Processing Systems*, Vol. 30 (2017).
- [49] M. Raissi, P. Perdikaris, and G. Karniadakis, Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations, *Journal of Computational Physics* **378**, 686 (2019).
- [50] C. Dankert, R. Cleve, J. Emerson, and E. Livine, Exact and approximate unitary 2-designs and their application to fidelity estimation, *Phys. Rev. A* **80**, 012304 (2009).
- [51] D. M. Debroy, É. Genois, J. A. Gross, W. Mruczkiewicz, K. Lee, S. Hong, Z. Chen, V. Smelyanskiy, and Z. Jiang, Context-aware fidelity estimation, *Phys. Rev. Res.* **5**, 043202 (2023).
- [52] More generally, a set of 4^n states are required to form a unitary 2-design on n qubits.
- [53] R. W. Heeres, P. Reinhold, N. Ofek, L. Frunzio, L. Jiang, M. H. Devoret, and R. J. Schoelkopf, Implementing a universal gate set on a logical qubit encoded in an oscillator, *Nature Communications* **8**, 10.1038/s41467-017-00045-1 (2017).
- [54] F. Mezzadri, How to generate random matrices from the classical compact groups, *Notices of the American Mathematical Society* **54**, 592 (2007).
- [55] S. Gustavsson, O. Zwiernik, J. Bylander, F. Yan, F. Yoshihara, Y. Nakamura, T. P. Orlando, and W. D. Oliver, Improving quantum gate fidelities by using a qubit to measure microwave pulse distortions, *Phys. Rev. Lett.* **110**, 040502 (2013).
- [56] M. A. Rol, L. Ciorciaro, F. K. Malinowski, B. M. Tarasinski, R. E. Sagastizabal, C. C. Bultink, Y. Salathe, N. Haandbaek, J. Sedivy, and L. DiCarlo, Time-domain characterization and correction of on-chip distortion of control pulses in a quantum processor, *Applied Physics Letters* **116**, 054001 (2020).
- [57] S. Lazăr, Q. Ficheux, J. Herrmann, A. Remm, N. Lacroix, C. Hellings, F. Swiadek, D. C. Zanuz, G. J. Norris, M. B. Panah, A. Flasby, M. Kerschbaum, J.-C. Besse, C. Eichler, and A. Wallraff, Calibration of drive nonlinearity for arbitrary-angle single-qubit gates using error amplification, *Phys. Rev. Appl.* **20**, 024036 (2023).
- [58] E. Hyyppä, A. Vepsäläinen, M. Papič, C. F. Chan, S. Inel, A. Landra, W. Liu, J. Luus, F. Marxer, C. Ockeloen-Korppi, S. Orbell, B. Tarasinski, and J. Heinsoo, Reducing leakage of single-qubit gates for superconducting quantum processors using analytical control pulse envelopes (2024), [arXiv:2402.17757](https://arxiv.org/abs/2402.17757).
- [59] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, T. Killeen, Z. Lin, N. Gimelshein, L. Antiga, A. Desmaison, A. Kopf, E. Yang, Z. DeVito, M. Raison, A. Tejani, S. Chilamkurthy, B. Steiner, L. Fang, J. Bai, and S. Chintala, Pytorch: An imperative style, high-performance deep learning library, in *Advances in Neural Information Processing Systems*, Vol. 32 (2019).
- [60] D. P. Kingma and J. Ba, Adam: A method for stochastic optimization (2017), [arXiv:1412.6980](https://arxiv.org/abs/1412.6980).
- [61] M. Paris and J. Rehacek, *Quantum State Estimation*, Lecture Notes in Physics (Springer-Verlag, Berlin Heidelberg, 2004).
- [62] M. A. Nielsen, A simple formula for the average gate fidelity of a quantum dynamical operation, *Physics Letters A* **303**, 249–252 (2002).

Appendix A: Data requirements

The success of our approach to quantum optimal control relies on constructing an accurate model of the quantum dynamics from data. We explore here how the amount and quality of training data impact our results. We change the accuracy of the data labels by numerically sampling more measurements (N_{shots}), from the experimentally inexpensive 32 shots (results of the main text), up to the limit of perfect labels. See Appendix C for a discussion on different approaches to sampling the pulse shapes we use for training.

In Fig. 5, we present the impact of the training dataset size and experiment repetitions (N_{shots}) used on both tasks considered in this work: the ML-based characterization of the quantum dynamics in panel (a) and the resulting quantum optimal control gate fidelity in panel (b). Unsurprisingly, using more training data leads to models with better prediction accuracy on the unseen test data. However, we show that such an improvement in the ML model is not necessarily needed to obtain high-fidelity quantum controls. For example, we can get the MSE on the model prediction down from 2.1×10^{-4} to

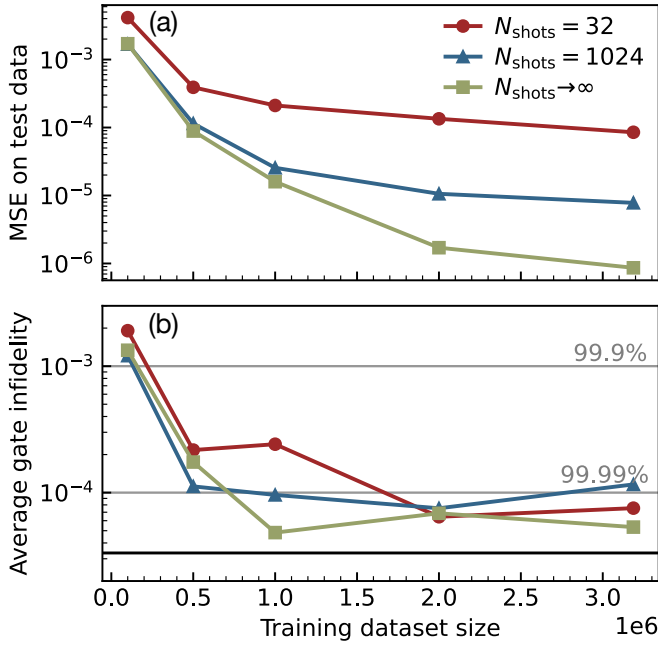


FIG. 5. Impact of the quantity and quality of training data on the model learning and optimal control performance. (a) Mean-squared prediction error on the unseen test dataset of RNN models trained on increasingly more pulse shapes with different sampling repetitions (N_{shots}). (b) Average gate infidelity of 20 ns $R_X(\pi)$ gate optimized on the models of panel (a), as evaluated on the true 5-level transmon model. The black line is the coherence limit of the gate.

8.5×10^{-5} by using 3.2 million training pulse shapes instead of 1 million. However, the optimal controls found by both of these models lead to a single-qubit gate fidelity of about 99.99%, with most of the variability in the resulting fidelities explained by the stochastic nature of both the model training and the control optimization. Similarly, increasing the number of shots leads to significant improvements in the model prediction accuracy (about an order of magnitude reduction in the MSE at 3.2 million pulses between the three cases considered), but this improvement does not translate into significant improvements in the resulting control fidelities.

We explain this feature of our approach from the fact that even though the ML model has a finite precision, which can be thought of as error bars on the model parameter estimates, as long as the learned parametrization does not lead to a significant bias in the quantum dynamics, the optimal controls on the finite precision model should also be close to optimal on the true quantum system. We can think of the QOC optimization on the trained neural network, which has a finite precision but very small bias, as adding Gaussian noise to the quantum state evolved using the master equation. The controls that maximize the gate fidelity on average using such a noisy model should also be optimal on the true noiseless model. Here, we show that as long as we have enough data to train a model with a mean-squared pre-

diction error of less than about 2×10^{-4} , we can obtain quantum gates with 99.99% fidelity. Given the important time cost associated with acquiring large amounts of data on the quantum system, it is then desirable to use as little data as required for the MLQOC protocol to yield the desired control fidelities.

We emphasize that on the MHz-scale repetition rate typical of quantum systems such as superconducting qubits, only a few hours are required to acquire the amount of data considered here. Given that our framework can be easily extended to parallel single-qubit gates with virtually no additional experimental time, optimizing quantum operations using MLQOC is easily achievable with current experiments. Additionally, the limited data requirements we obtain for performing arbitrary single-qubit operations open the door to realistically using MLQOC for more complex control tasks such as two-qubit gates, readout, and reset operations, which is beyond the scope of this work.

Appendix B: Machine-learning trainings

In this section, we present the implementation details of our machine-learning (ML) model trainings, whose architecture is illustrated in Fig. 2. Our numerical implementation is based on PyTorch [59]. The recurrent neural network we use is the vanilla long-short term memory unit implemented in `torch.nn.lstm` with a single layer and a hidden size of 48, see the PyTorch documentation for more details. The *Encode* transformation is depth-2 fully connected neural networks (FCNN) with a hidden size of 96 and a sigmoid activation function, whereas the *Decode* transformation is a single-layer linear network, using again a sigmoid activation function to map the network output into proper probabilities $\in (0, 1)$. We emphasize that none of these hyperparameters were carefully optimized, as the objective in this work is to demonstrate how our proposed MLQOC framework can be successful with a simple implementation. Such an hyperparameter optimization, together with using more powerful neural network architectures such as transformer models [48], could yield slightly better performances, but would not change any of our conclusions.

We train this model using gradient-descent based on the Adam optimizer [60] with a learning rate of 0.001. During training, we use an initial batch size of 256 and double it at epochs [100, 200, 500, 800] in order to reduce the stochastic noise in the optimal model parameters as training progresses. We use the following physics-

inspired loss function to train our model [24]

$$\mathcal{L} = \mathcal{L}_{\text{pred}} + w_{\text{posit}} \mathcal{L}_{\text{posit}} + w_{\text{prep}} \mathcal{L}_{\text{prep}}, \quad (\text{B1})$$

$$\mathcal{L}_{\text{pred}} = \frac{1}{N} \sum_{n=1}^N \left((\Pi_n^j - \tilde{\Pi}_n^j)^2 \right), \quad (\text{B2})$$

$$\mathcal{L}_{\text{posit}} = \frac{1}{N(N_t + 1)} \sum_{t=0}^{N_t} \sum_{n=1}^N \text{ReLU} \left(|\tilde{\rho}_{t,n}|^2 - 1 \right), \quad (\text{B3})$$

$$\mathcal{L}_{\text{prep}} = \frac{1}{N} \sum_{n=1}^N |\tilde{\rho}_{0,n} - \rho_{0,n}|^2, \quad (\text{B4})$$

where $\tilde{\rho}_{t,n} = (\tilde{x}_{t,n}, \tilde{y}_{t,n}, \tilde{z}_{t,n})$ is the model prediction of the qubit states for the n th quantum trajectory at time index t , given N pulse shapes with N_t pixels. We use the hyperparameters $w_{\text{posit}} = w_{\text{prep}} = 1.0$.

The prediction loss $\mathcal{L}_{\text{pred}}$ is a simple mean squared error between the labels and the model predictions for the qubit population in the $j \in \{X, Y, Z\}$ basis, given by $\tilde{\Pi}_n^j = (1 - \langle \sigma^j \rangle) / 2 \in [0, 1]$. Here the labels are also probabilities of measuring the qubit state in the -1 eigenstate of the operator σ^j , which are constructed from averaging the $N_{\text{shots}} = 32$ projective measurements for each input pulse shape. The MSE given by this loss term is the metric we care most about to obtain an accurate quantum characterization of the qubit dynamics. However, there is more information about the RNN output that we can use to quantify if it accurately describes our physical problem.

We leverage that information with the following two loss terms, which allow to make our model training more accurate and more data efficient. The positivity loss $\mathcal{L}_{\text{posit}}$ insures that the ML model predictions correspond to valid quantum states, i.e. that the predicted density matrices are positive. The rectified linear unit $\text{ReLU}(x) = \max(0, x)$ allows us to penalize only for non-physical states and not for mixed state predictions, given that we are learning from data with relaxation and dephasing. Finally, the preparation loss $\mathcal{L}_{\text{prep}}$ makes the RNN predictions accurate at the beginning of the quantum trajectory, $t = 0$, by using our knowledge of the finite set of prepared initial states. When dealing with experimental data, or more generally with data containing state preparation and measurement (SPAM) errors, we can set the targeted prepared states $\rho_{0,n}$ to be consistent with our best estimate of these errors. For example, we can average over subsets of the data where the projective readout immediately follows the preparation, and effectively perform quantum state tomography of the six cardinal states of the Bloch sphere [61].

We note that the model architecture, together with the loss functions we use, can easily be extended to modeling the dynamics of n qubits or multi-level systems. For example, the input dimension of the RNN could be $2n$ for n microwave drives with two quadratures, and the prepared and output quantum states can be represented using $4^n - 1$ expectation values for n qubits. This natural extension would be sufficient to apply our frame-

work to optimize for two-qubit gates and parallel single-qubit gates. Of course, the data requirements and model complexity necessary for successfully applying MLQOC to controlling a quantum system with an exponentially large Hilbert space should also scale unfavorably with the number of qubits. However, using the heuristic representation of a machine-learning model might be beneficial for tackling such complex control problems.

Appendix C: Pulse datasets

The quantum characterization task is achieved by training a parametrized representation, such as a neural network acting as a universal function approximator, to describe the transformation from arbitrary input pulse shapes into the dynamics of the quantum systems under study. Given the enormous size of the input control parameter space, with every pixel taking an arbitrary floating point value, and the fact that we want to acquire a finite dataset in a reasonable experimental time, it is natural to try uniformly sampling at random the input controls to form our ML datasets. This approach works well and the trained ML model can predict the dynamics for random pulse shapes accurately. However, with that choice the model performs poorly at predicting the dynamics of smooth pulses, which for example possess a very different frequency spectrum than random pulses (data not shown). Given that in an experimental scenario we want to use smooth control pulses, such as to respect the bandwidth limitations of the control electronics and avoid crosstalk problems arising from having signals with a wide frequency spectrum, having a trained model that is inaccurate at capturing the dynamics of smooth pulses is practically not useful.

To remedy the situation, we simply include smooth pulses in the training data. Here, we generate these smooth pulses using physically-motivated envelopes that are commonly used to perform single-qubit gates. Specifically, in addition to uniformly sampled random envelopes, our set of pulses is composed of flat, gaussian, gaussian with a flat top, gaussian with an orthogonal DRAG component, and sinusoidal pulses. Heuristically, we chose proportions of 25 % of flat pulses with random amplitudes, 25 % of flat-top gaussian pulses with random widths, amplitudes, and standard deviations, 25 % of uniformly sampled random pulses, 12.5 % of gaussian pulses with an orthogonal DRAG component with random amplitudes, and 12.5 % of sinusoidal pulses with random amplitudes and frequencies. We generated a total 4.25 million pulses, and used 75 % for the training set, 15 % for the validation set used to avoid overfitting, and 10 % for the test set used to evaluate and compare model performances.

Given our pulse parametrization at the level of the AWG before IQ mixing, our pulses should have a main sinusoidal oscillation corresponding to the intermediate frequency, here around 100 MHz, such that the resulting

pulses reach the qubit frequency and produce non-trivial dynamics. We thus generate all of the aforementioned pulses at the envelope level, i.e. without any explicit oscillating component. We then convolve these envelopes with the intermediate frequency oscillations, which result in the final pulses of the dataset. In order to explicitly sample the impact of detuned drives on the qubit dynamics, we randomly sample the intermediate frequency we use to convolve with our envelopes, using a gaussian distribution centered at the nominal $\omega_{\text{IF}}/2\pi = 100$ MHz with a standard deviation of 1 MHz.

We note that fine tuning the quality of the training dataset by using pulses that are known to achieve the target operations with good fidelity, and perhaps increasing the sampling (N_{shots}) for these pulses, could be beneficial to the performance of the MLQOC approach. However, as we have demonstrated in the main text, such fine tuning is unnecessary and the approach works well as long as the pulses seen during training have the same characteristics as the controls we implement in the end. In that sense, although such an informed construction of the training dataset will necessarily bias the model learning towards a given set of dynamics, this bias should be beneficial for learning a good model efficiently, as long as the bias is consistent with the constraint we put on the following QOC optimizations. Indeed, it is desirable to have a trained model that is good at predicting the dynamics of smooth pulses, even if it is somewhat specialized and not fully generalizable, since ultimately this is what our optimal control task requires. This bias is enforced explicitly in this work by using a cost function minimizing the first and second derivatives of the optimal controls.

Appendix D: QOC implementation details

In this section, we present the implementation details of our quantum optimal control (QOC) approach illustrated in Fig. 1(b). We use an open-loop optimization where the pulses are parametrized at every pixel, the trained ML model with fixed parameters is used to compute the dynamics resulting from the pulses, and the cost

function is expressed as follows [25]

$$\mathcal{C} = \mathcal{C}_{\text{fidel}} + w_{\text{clamp}}\mathcal{C}_{\text{clamp}} + w_{\text{mean}}\mathcal{C}_{\text{mean}} \quad (\text{D1})$$

$$+ w_{\text{first}}\mathcal{C}_{\text{first}} + w_{\text{second}}\mathcal{C}_{\text{second}}, \quad (\text{D2})$$

$$\mathcal{C}_{\text{fidel}} = \frac{1}{N} \sum_{n=1}^N 1 - \text{AGF}(\rho_{n,T}), \quad (\text{D3})$$

$$\mathcal{C}_{\text{clamp}} = \frac{1}{N(N_t + 1)} \sum_{t=0}^{N_t} \sum_{n=1}^N \text{ReLU}(\Omega_{n,t} - \Omega_{\text{max}}), \quad (\text{D4})$$

$$\mathcal{C}_{\text{mean}} = \frac{1}{N(N_t + 1)} \sum_{t=0}^{N_t} \sum_{n=1}^N |\Omega_{n,t}|^2, \quad (\text{D5})$$

$$\mathcal{C}_{\text{first}} = \frac{1}{NN_t} \sum_{t=1}^{N_t} \sum_{n=1}^N |\partial_t \Omega_{n,t}|^2, \quad (\text{D6})$$

$$\mathcal{C}_{\text{second}} = \frac{1}{N(N_t - 1)} \sum_{t=2}^{N_t} \sum_{n=1}^N |\partial_t^2 \Omega_{n,t}|^2, \quad (\text{D7})$$

where Ω_{max} is the maximal allowed drive amplitude, which is $2\pi \times 100$ MHz in this work. This constitutes a reasonable choice to avoid significantly amplifying the effect of classical crosstalk on a transmon chip, as typical sinusoidal 20 ns single-qubit π gates require $|\Omega|/2\pi \approx 31$ MHz. The partial time derivatives in the smoothing cost functions are implemented as finite differences numerically with $\partial_t \Omega_{n,t} = \Omega_{n,t} - \Omega_{n,t-1}$, and the average gate fidelity is computed as [62]

$$\text{AGF}(\rho_{n,T}) = \frac{\sum_j \text{Tr} \left(U_{\text{target}} \rho^j U_{\text{target}}^\dagger \text{RNN}(\rho^j) \right) + d^2}{d^2(d+1)}, \quad (\text{D8})$$

where $d = 2$ for a qubit, T the final time index, σ^j are the 6 cardinal states (our chosen unitary 2-design), and $\text{RNN}(\rho^j)$ represents the transformation realized by our trained ML model, which acts here in place of the usual quantum channel.

As mentioned in the main text, we optimize for a batch of 30 pulses in parallel on the same GPU card in under a minute, and select the best performing pulse on the full cost function as the optimal control. We use the Adam optimizer with a learning rate of 0.001 to perform the gradient descent [60]. As in the ML model training, all the gradients are computed numerically using the auto-differentiable feature of PyTorch [59].

We have observed that whereas the clamping cost $\mathcal{C}_{\text{clamp}}$ simply allows the pulses to respect an experimental constraint, the average amplitude cost $\mathcal{C}_{\text{mean}}$ together with the smoothing cost functions $\mathcal{C}_{\text{first}}$ and $\mathcal{C}_{\text{second}}$ are very important for constraining the parameter exploration during the optimization towards controls where the trained RNN has an accurate representation of the dynamics. Indeed, the trained model predicts dynamics which are very close to the true master equation dynamics for smooth input pulses, which allows us to find optimal controls with over 99.99% fidelity when using these

regularizing cost functions. However, optimizing the controls on the trained model using only the fidelity cost can lead to pulses where the predicted and true dynamics diverge significantly. This problem can be completely avoided using the smoothing cost functions, and if necessary a simple post-selection based on the same smoothness criteria over the optimization results of the batch. We optimize our pulses for all the single-qubit gate results presented in the main text using $w_{\text{clamp}} = 10.0$, $w_{\text{mean}} = 0.1$, and $w_{\text{first}} = w_{\text{second}} = 0.01$. We have observed that fine tuning these hyperparameters do not lead to significant changes in the resulting optimal gate fidelities.

Appendix E: Random pulse distortions

We simulate pulse distortions numerically by applying a fixed causal transfer function to the pulses. We first generate the random transfer function $F(\omega)$ in the frequency domain, using the noise parameter θ that we scale from 0 to 0.4 in Fig. 4. We add the noise to the trivial flat transfer function such that

$$F(\omega \geq 0) = 1 + U(-\theta, \theta), \quad (\text{E1})$$

where $U(\min, \max)$ is the uniform distribution, and we replicate the values for the negative frequencies such that $F(\omega)$ is even. To obtain a smooth transfer function, we sample 11 points from the uniform distribution, that we attribute to frequencies up to 600 MHz, before using a gaussian filter with a 2.5 GHz bandwidth to interpolate between these values for 1001 points in the range $[-1.5, 1.5]$ GHz.

Instead of using Fourier transforms to convolve all the pulses of our datasets with the noisy transfer function defined in the frequency domain, we express $F(\omega)$ in the time domain using a transfer matrix which can be directly applied to the pulses using a simple matrix multiplication. This transfer matrix is obtained by numerically solving [7]

$$T_{jk} = \int_{-\infty}^{\infty} F(\omega) \frac{\sin(\omega \Delta t / 2) \cos(\omega(j-k)\Delta t)}{\pi \omega} d\omega, \quad (\text{E2})$$

where Δt is the pixel size (in units of time) and j, k are the pixel indices.

Finally, given that this noise models a physical process, we impose causality by setting the lower triangular elements to zero and renormalizing it. Using this approach, we can simulate a realistic scenario where out-of-model dynamics are present in the quantum system we are trying to control, and we can have a single tunable parameter θ to quantify the resulting model bias.

Appendix F: IQ mixing and qubit drive

In this section, we detail the IQ mixing transformation applied to the input microwave pulses, typically gener-

ated by an arbitrary waveform generator (AWG), before reaching the transmon qubit. We then show how these mixed pulses can be viewed as the usual σ_X and σ_Y drive Hamiltonian prefactors. As detailed in Refs. [41, 42], the two AWG signals S_I and S_Q get mixed with a local oscillator signal of frequency ω_{LO} in a process known as sideband mixing to produce the drive

$$\Omega(t) = S_I(t) \cos(\omega_{\text{LO}} t) + S_Q(t) \sin(\omega_{\text{LO}} t). \quad (\text{F1})$$

We then define the AWG signals with an envelope combined with sinusoidal oscillations at an intermediate frequency ω_{IF} and a phase ϕ , such that

$$S_I(t) = I(t) \cos(\omega_{\text{IF}} t + \phi), \quad (\text{F2})$$

$$S_Q(t) = -Q(t) \sin(\omega_{\text{IF}} t + \phi). \quad (\text{F3})$$

We then directly get

$$\Omega(t) = \frac{I(t)}{2} [\cos(\omega_+ t + \phi) + \cos(\omega_- t + \phi)] \quad (\text{F4})$$

$$- \frac{Q(t)}{2} [\cos(\omega_- t + \phi) - \cos(\omega_+ t + \phi)] \quad (\text{F5})$$

where $\omega_{\pm} = \omega_{\text{IF}} \pm \omega_{\text{LO}}$. We thus see that by setting $I(t) = Q(t) = A(t)$, we obtain a single carrier frequency signal at the upper sideband ω_+ ,

$$\Omega(t) = A \cos(\omega_+ t + \phi). \quad (\text{F6})$$

Using a LO close to the qubit frequency such that $\omega_{\text{IF}} + \omega_{\text{LO}} \approx \omega_q$, we can then precisely control the frequency spectrum of the qubit pulses we send, with a limitation set by the AWG bandwidth and potentially the cable filters used in the experiment. Note that one can also choose a different phase between the S_I and S_Q signals to drive the lower sideband.

We now demonstrate that such a drive can be used to perform arbitrary qubit gates. Starting from Eq. (2), we can introduce the creation and annihilation operators to diagonalize the quadratic terms of the transmon Hamiltonian and obtain, after expanding the $\cos \hat{\phi}$ term and truncating after the second order [39]

$$\hat{H}(t) \approx \omega_q \hat{b}^\dagger \hat{b} - \frac{E_C}{2} \hat{b}^\dagger \hat{b}^\dagger \hat{b} \hat{b} + \Omega(t) \frac{i}{2} \left(\frac{2E_C}{E_J} \right)^{1/4} (\hat{b}^\dagger - \hat{b}). \quad (\text{F7})$$

Going into the qubit rotating frame with the transformation $\hat{U}(t) = \exp(i\omega_q \hat{b}^\dagger \hat{b} t)$ and applying the rotating-wave approximation (RWA), we obtain the time-dependent drive Hamiltonian

$$\hat{H}_d(t) \approx \frac{i}{2} \left(\frac{2E_C}{E_J} \right)^{1/4} \frac{A(t)}{2} \left(\hat{b}^\dagger e^{i(\Delta t + \phi)} - \hat{b} e^{-i(\Delta t + \phi)} \right), \quad (\text{F8})$$

where $\Delta = \omega_+ - \omega_q$. Performing a two-level truncation with $\hat{b}^\dagger \rightarrow \sigma_+$, $\hat{b} \rightarrow \sigma_-$, and $\sigma_{\pm} = (\sigma_X \pm i\sigma_Y)/2$, redefining

the prefactor in front of the last parenthesis to be $\tilde{A}(t)$ and assuming that the drive is on resonance with the qubit ($\Delta = 0$), we finally get the textbook single-qubit drive Hamiltonian

$$\hat{H}_d^{\text{qubit}}(t) = \tilde{A}(t) [\cos(\phi)\sigma_X - \sin(\phi)\sigma_Y]. \quad (\text{F9})$$

We thus understand how controlling the pixel amplitudes of the AWG inputs $S_I(t)$ and $S_Q(t)$ allows us to perform arbitrary single-qubit operations.

4.5 Perspectives

Alors que ce projet se concentre sur l'introduction de la méthodologie combinant une caractérisation par apprentissage automatique avec le contrôle optimal quantique ainsi que la démonstration de son fonctionnement et de son utilité, de nombreuses avenues de recherche intéressantes restent inexplorées. En plus d'une démonstration expérimentale, je discute ici de trois directions principales que je considère comme particulièrement prometteuses pour faire suite au projet : l'étude du compromis biais-variance avec l'entraînement de modèles d'apprentissage différents, l'optimisation d'opérations quantiques plus complexes, ainsi que l'application directe du modèle entraîné pour améliorer notre compréhension et notre utilisation du dispositif quantique.

4.5.1 Modularité et compromis biais-variance

L'un des aspects rendant notre approche particulièrement intéressante pour résoudre divers problèmes de contrôle quantique est sa modularité. En effet, les différentes sous-composantes de notre méthodologie, telles que l'expérience spécifique réalisée pour générer les données d'entraînement, l'architecture d'apprentissage, ainsi que l'algorithme de contrôle optimal, peuvent toutes être choisies judicieusement afin de satisfaire les besoins spécifiques en efficacité expérimentale et en performance du dispositif quantique étudié. Par efficacité expérimentale, j'entends simplement la durée totale et le nombre d'expériences différentes qu'il est possible de réaliser afin de générer les données d'entraînement et d'obtenir les solutions de contrôle recherchées. Par exemple, la dérive des paramètres expérimentaux peut faire en sorte qu'une opération calibrée ne possède une grande fidélité que pour une journée. Combinée à la fréquence de répétition possible de l'expérience, qui est par exemple de l'ordre de ~ 100 kHz pour des qubits supraconducteurs et de ~ 10 Hz pour des ions piégés [49], cette contrainte peut mener à une limite supérieure restrictive sur la taille de l'ensemble de données accessible. Étant donné que chaque réalisation expérimentale possède ses contraintes et ses particularités qu'il est possible d'exploiter, l'aspect modulaire de l'approche utilisée devient nécessaire afin d'assurer son bon fonctionnement.

Plus spécifiquement, alors que la performance des contrôles obtenus par notre approche MLQOC est limitée par la capacité du modèle entraîné à représenter les dynamiques du système, le choix d'architecture d'apprentissage s'avère déterminant. Alors que nous nous sommes concentrés sur l'utilisation de réseaux de neurones récurrents dans notre preuve de concept, ce choix n'est pas nécessairement optimal étant donné la situation expérimentale rencontrée. On sait par exemple que dans le cadre de nos simulations, le modèle réel décrivant

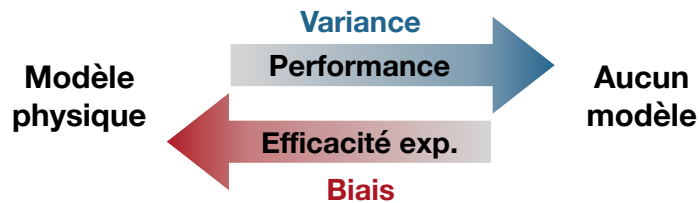


FIGURE 4.6 – Compromis entre biais et variance du modèle d’apprentissage utilisé. L’utilisation d’un modèle physique permet l’utilisation d’une transformation paramétrée du problème qui est simplifiée et qui nécessite donc un ensemble minimal de données pour être entraînée. Cependant, l’expressivité restreinte de ce modèle, lequel n’est en mesure de représenter que certaines transformations bien définies, limite sa performance. De l’autre côté, s’affranchir de tout modèle permet d’apprendre une transformation complètement générale et donc de tenir compte de l’ensemble des corrélations présentes dans les données, menant ainsi à une performance optimale en présence d’un ensemble suffisant de données d’entraînement. Le désavantage de cette approche est qu’elle est très peu efficace expérimentalement, nécessitant de nombreuses expériences pour apprendre à partir de zéro.

les dynamiques du transmon, décrit par l’équation maîtresse choisie, permet d’atteindre des fidélités de porte significativement plus élevées, lesquelles sont ultimement limitées par la cohérence du transmon (ici 99.997 % en comparaison au 99.98 % atteint avec notre approche). Comme discuté un peu plus bas, il devrait alors être possible d’utiliser la structure de l’équation maîtresse afin d’entraîner un modèle plus précis et ainsi atteindre de meilleures opérations quantiques. Cependant, l’utilisation de telles suppositions sur la dynamique du système nous expose au problème de biais du modèle qui peut à son tour limiter la performance des contrôles en pratique si certaines dynamiques qui ont été supposées ne sont pas présentes dans le dispositif expérimental.

Ce compromis dans le choix du modèle d’apprentissage est illustré à la Fig. 4.6, considérant à la fois la capacité du modèle à décrire les dynamiques quantiques (performance), ainsi que la quantité de données nécessaire pour y arriver (efficacité expérimentale). La recherche d’un bon équilibre entre l’expressivité et le biais de modèle est cruciale parce qu’en pratique, on ne possède toujours qu’un ensemble fini de données bruitées. La sélection de la bonne architecture d’apprentissage permet alors de tirer le maximum de ces données pour son application en contrôle optimal.

L’intuition rendant cette problématique de recherche intéressante peut être décrite comme suit. D’un côté, une approche n’utilisant aucun modèle telle que celle activement poursuivie en apprentissage par renforcement peut rapidement s’avérer beaucoup trop coûteuse en temps ou en données pour de nombreuses applications de contrôle quantique. De l’autre côté, l’utilisation d’un modèle basé uniquement sur des principes physiques ne permet typiquement pas de décrire l’ensemble des dynamiques avec le niveau de précision nécessaire. Ainsi, il devient avantageux de chercher le juste milieu dans ce dilemme biais-

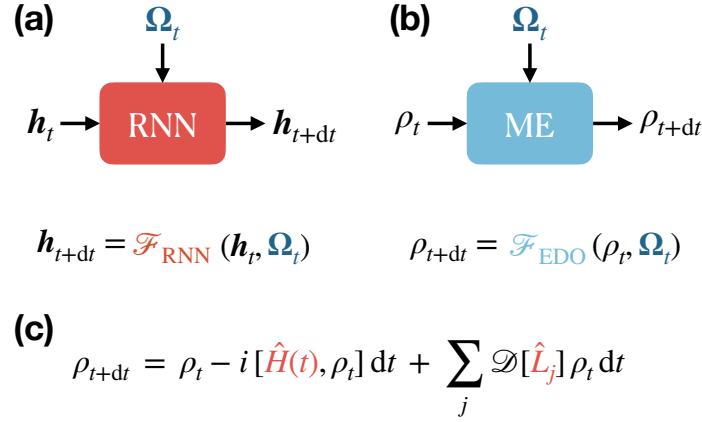


FIGURE 4.7 – Apprendre les dynamiques quantiques directement à partir d’une représentation paramétrée de l’équation maîtresse. (a) Le réseau de neurones récurrent (RNN) utilise un état caché (h_t) pour décrire les dynamiques du dispositif quantique. (b) Le modèle physique d’une équation maîtresse (ME) représente l’état du système à l’aide d’une matrice densité (ρ_t) et modélise son évolution temporelle à l’aide d’une équation différentielle ordinaire (EDO), laquelle réalise une transformation complètement positive tout en préservant la trace de la matrice densité. (c) En utilisant des paramètres variationnels pour représenter certaines portions de cette transformation, par exemple l’Hamiltonien et les opérateurs de dissipation accentués en orange, il est possible d’utiliser direction l’équation maîtresse comme modèle d’apprentissage. Ce genre d’approche est grandement facilité par l’utilisation de la librairie `dynamiqs`, disponible en libre accès [79].

variance afin d’entraîner de façon efficace le modèle nous permettant d’atteindre le niveau de contrôle quantique recherché. Ce genre d’approche, parfois qualifiée de boîte grise [160], ouvre un continuum de possibilités entre l’utilisation de modèles purement physiques (par exemple avec un Hamiltonien et une équation maîtresse) et de modèles qui sont des boîtes noires (tel qu’un réseau de neurones).

Afin de rendre le tout un peu plus concret, je présente le lien explicite entre l’utilisation d’un réseau de neurones récurrent et d’une équation maîtresse pour décrire les dynamiques d’un système quantique à la Fig. 4.7, en plus de présenter comment une telle équation peut être utilisée comme modèle d’apprentissage. Le RNN apprend à suivre l’évolution de l’état quantique du système de façon implicite à l’aide d’un état caché (h_t) qu’il modifie en fonction des contrôles reçus en entrée de façon séquentielle. D’une manière similaire, l’équation maîtresse transforme un état quantique (ρ_t) à chaque pas de temps en fonction des contrôles externes venant modifier l’Hamiltonien du système.

L’idée d’utiliser un modèle physique afin de caractériser le dispositif quantique est réalisée en définissant certains éléments de l’équation maîtresse comme paramètres variables. Le modèle d’apprentissage devient alors une équation différentielle ordinaire (EDO) paramétrée et que l’on peut résoudre à l’aide de l’un des nombreux algorithmes numériques

existants, tels que les méthodes de Runge-Kutta [46]. Je note que ce modèle d'équation maîtresse paramétrée peut être entraîné de manière très efficace sur des processeurs graphiques en utilisant une librairie supportant l'autodifférentiation, telle que PyTorch ou JAX [78]. Une implémentation de la sorte est d'ailleurs grandement facilitée par la librairie en libre accès dynamiqs sur laquelle j'ai contribué [79]. Je décrirai cette librairie et ses applications en contrôle optimal et en estimation de paramètre au chapitre 6.

Tel qu'illustré à la Fig. 4.7(c), il est possible de choisir comme paramètres libres les entrées d'une matrice hermitienne représentant l'Hamiltonien du système ainsi qu'un certain nombre de matrices unitaires représentant les opérateurs de dissipation. Cela constitue une modélisation purement physique, ou sous forme de boîte blanche, du problème. Cependant, ce choix est loin d'être unique et il est possible de tester de nombreux modèles physiques comprenant plus ou moins de degrés de libertés ou de sources de bruits afin de quantifier les dynamiques présentes dans le dispositif quantique. Une telle caractérisation par sélection de modèles [161] promet d'être très utile pour acquérir une compréhension approfondie du système quantique que l'on désire contrôler. Au-delà des différentes modélisations physiques, il est également possible d'ajouter des paramètres arbitraires au modèle d'apprentissage afin de tenir compte des dynamiques qui seraient trop coûteuses à simuler rigoureusement, comme celles associées aux éléments supraconducteurs dont le qubit est couplé capacitivement et inductivement comme ses (quatre) coupleurs, sa cavité de mesure, les défauts à deux niveaux dans les matériaux l'entourant, etc. L'obtention de modèles heuristiques précis pour décrire ces dynamiques demeure un objectif de recherche important dans le développement des technologies quantiques.

4.5.2 Opérations quantiques plus complexes

Au-delà des opérations logiques à un qubit démontrées dans ces travaux, il serait fort intéressant d'explorer les performances de notre approche pour réaliser des opérations plus complexes impliquant davantage de qubits. De par sa simplicité et sa modularité, l'approche MLQOC pourrait en effet s'avérer bénéfique afin d'améliorer plusieurs opérations logiques complexes réalisées en pratique, lesquelles ont typiquement des fidélités plus faibles que celle de 99.99 % aujourd'hui accessible pour des portes à un qubit en isolation.

La première extension naturelle mais non triviale consiste à considérer les portes à un qubit réalisées en parallèle. Sur le système à Berkeley et de nombreux autres, ces opérations possèdent une erreur moyenne environ dix fois plus grande que lorsque réalisées en isolation [162]. Cette réduction de notre capacité à contrôler précisément le dispositif quantique est principalement due à la diaphonie classique et quantique. Par diaphonie classique,

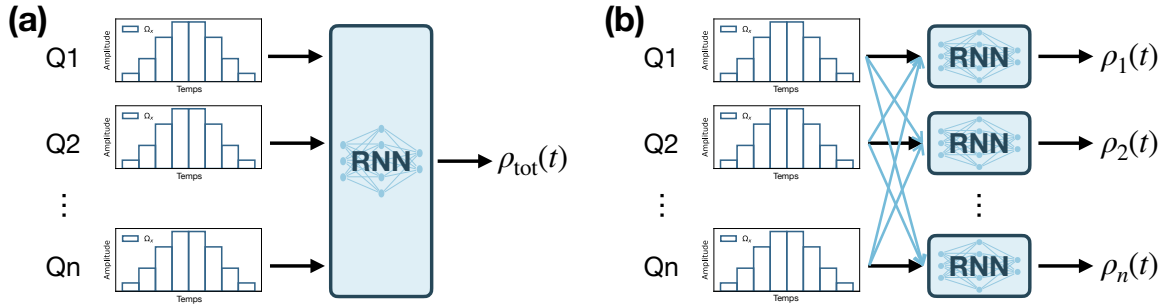


FIGURE 4.8 – Optimisation de portes parallèles à un qubit à l’aide de l’approche MLQOC. (a) L’utilisation d’un modèle représentant l’ensemble des dynamiques temporelles de n qubits est l’approche la plus générale, mais nécessite d’apprendre un nombre exponentiel de degrés de liberté. (b) De façon alternative, le problème d’apprentissage peut être décomposé en différentes portions de manière à tenir compte des dynamiques principales en jeu dans le système.

j’entends simplement qu’en raison des différents couplages capacitifs de l’échantillon, le signal alimentant un qubit peut également affecter d’autres qubits avec des atténuations et des délais temporels différents. Par diaphonie quantique, je fais référence aux interactions générant de l’enchevêtrement entre les différentes composantes du système, telles que l’interaction capacitive entre deux qubits produisant un couplage ZZ [64]. La présence de ce couplage non nul déplace la fréquence de résonance du qubit de façon conditionnelle à l’état de nombreux autres qubits. De façon formelle, cet effet modifie le contrôle individuel des n qubits en parallèle ($n \times U_{2 \times 2}$) vers un problème de contrôle sur n qubits ($U_{2^n \times 2^n}$).

Le modèle d’apprentissage à utiliser afin d’appliquer notre approche MLQOC au problème d’opérations logiques parallèles devrait idéalement tenir compte de tous les degrés de liberté du système. Tel qu’illustré à la Fig. 4.8(a), le modèle serait alors en mesure de prédire l’évolution du système total à partir de contrôles arbitraires sur chaque qubit. Étant donné que ces degrés de liberté augmentent de façon exponentielle avec la taille du système (2^n), on peut s’attendre à ce que la complexité du modèle d’apprentissage de même que le nombre de données nécessaires pour l’entraîner soient très élevés, même pour des processeurs ne comportant que quelques qubits.

Une approche beaucoup plus réaliste consiste à décomposer l’architecture du modèle d’apprentissage en cellules unitaires, où chaque sous-modèle pourrait par exemple apprendre à décrire la dynamique d’un seul qubit. Cet exemple est illustré à la Fig. 4.8(b). Cette approche se base sur le fait que la diaphonie quantique impacte les dynamiques de façon beaucoup plus faible que les impulsions micro-ondes utilisées, ce qui est le cas pour la grande majorité des réalisations physiques. En effet, les qubits voisins sont typiquement décalés dans le domaine fréquentiel afin de minimiser leurs interactions [163] ou des éléments ajustables de couplage sont utilisés pour minimiser l’interaction ZZ [164]. Afin de tenir

compte de la diaphonie classique, il est possible d'ajouter une certaine connectivité entre les nœuds (illustrée avec les flèches bleues) et d'utiliser notre connaissance des particularités du dispositif afin de choisir quelles connexions ajouter, par exemple sur la base de l'allocation fréquentielle des différents qubits. Il serait très intéressant de tester une telle approche de MLQOC et de la mettre à l'échelle sur un dispositif quantique comportant plusieurs qubits.

Dans le même ordre d'idées, il serait possible d'appliquer l'approche MLQOC aux portes à deux qubits générant de l'intrication. Ces opérations sont plus bruyantes avec des fidélités d'environ 99.5 % et pourraient être significativement accélérées à l'aide d'une approche de conception des impulsions par contrôle optimal [165]. La méthodologie développée ici est directement applicable à ce problème de contrôle alors que les étiquettes provenant de mesures dispersives des qubits individuels peuvent également être utilisées pour estimer les valeurs attendues de l'état à deux qubits. De plus, les entrées prennent la même forme d'impulsions micro-ondes pour les opérations à résonance croisée réalisant un CNOT [166], ou deviennent des contrôles de basse fréquence du flux externe pour les opérations conservant le nombre d'excitations comme le CZ [71, 101]. Dans les deux cas, le fonctionnement du modèle d'apprentissage reste inchangé, c'est-à-dire que l'on cherche à prédire les observables quantiques à partir des contrôles fournis en entrée.

Outre l'augmentation des espaces d'Hilbert et de contrôle, une particularité rendant cette tâche plus difficile provient de l'estimation des paramètres d'enchevêtrement résiduels, lesquels sont typiquement difficiles à estimer. Par exemple, afin de réaliser l'opération unitaire correspondant à un CZ, l'échange d'excitation entre les deux qubits, une transformation associée à un opérateur de la forme $XX + iYY$, doit être nulle. Il est alors nécessaire d'optimiser les contrôles utilisés afin de minimiser cette quantité qui est nécessairement non nulle en pratique. De par sa faible valeur comparable à de nombreuses autres sources de bruit, cette quantité est difficile à évaluer à partir d'un nombre restreint d'échantillons. Heureusement, il est possible de concevoir judicieusement les expériences produisant les données d'entraînement de sorte à obtenir une certaine visibilité sur ce type de petits paramètres que l'on cherche à contrôler, notamment en utilisant un principe d'amplification où l'opération étudiée est répétée un nombre variable de fois. Un exemple très pertinent d'une telle approche est l'utilisation du protocole MEADD [4] que je présenterai au prochain chapitre.

Enfin, je tiens à souligner qu'en plus des opérations unitaires, notre approche MLQOC peut être appliquée à l'optimisation d'opérations dissipatives telles que la mesure et la réinitialisation d'un qubit supraconducteur. Comme nous le verrons au chapitre 6, ces opérations peuvent significativement bénéficier d'une approche de contrôle optimal quantique. L'utilisation d'un modèle précis devient alors cruciale à la réalisation réussie

d'une telle approche en pratique. Pour le problème d'optimisation de la mesure du qubit, l'une des principales questions ouvertes concerne le type d'étiquette idéal à utiliser pour entraîner le modèle décrivant le système. Plus spécifiquement, les signaux continus produits par la mesure peuvent être utilisés de plusieurs façons, notamment en les intégrant directement dans la résolution d'une équation maîtresse stochastique [36], en calculant les corrélations à n points des distributions de signaux produits [167], en les moyennant et en intégrant une équation maîtresse [13], ou en utilisant une approche agnostique à la physique du système. Il serait très intéressant d'étudier laquelle de ces approches est la plus efficace et performante pour réaliser le contrôle optimal de la mesure d'un qubit supraconducteur.

4.5.3 Applications directes

Afin de compléter le survol des perspectives de recherche futures pour ce projet, je détaille ici trois applications directes associées au modèle d'apprentissage que l'on entraîne avec l'approche MLQOC. Ces applications sont illustrées à la Fig. 4.9 et concernent l'étalonnage du dispositif à plus grande échelle, le prototypage de nouveaux protocoles, ainsi que le contrôle par rétroaction.

Tel qu'illustré à la Fig. 4.9(a), l'étalonnage des différentes opérations réalisées dans un ordinateur quantique peut être représenté par un graphe orienté acyclique [88]. Notre approche de contrôle optimal pourrait être bénéfique dans plusieurs de ces nœuds où la tâche est également de trouver la valeur des contrôles réalisant un objectif quantifiable, tel que la minimisation d'un couplage entre deux qubits ou d'une population dans un état donné. De manière plus générale, chacune de ces calibrations consiste à envoyer des signaux de contrôle au système quantique et à recevoir des bits de mesure en retour. De ce fait, l'ensemble des données générées lors de l'étalonnage d'un dispositif pourrait potentiellement être utilisé pour entraîner le modèle d'apprentissage.

Étant donné que le modèle d'apprentissage construit une transformation approchant celle réalisée par le système quantique réel, ce modèle peut être utilisé en tant que *jumeau digital* [168] du dispositif. L'utilisation du modèle entraîné dans diverses tâches de simulations est bénéfique pour plusieurs raisons. Tout d'abord, la transformation des entrées (contrôles) vers les sorties (observables) s'évalue rapidement de façon numérique et est autodifférentiable. De ce fait, plusieurs balayages de paramètres nécessitant de nombreuses expériences pourraient être remplacés ou voir leur ampleur diminuée en utilisant une approche de descente du gradient avec le modèle. De plus, le modèle représente de façon heuristique un grand nombre d'imperfections pertinentes sur le plan expérimental qui ne sont typiquement pas prises en compte dans notre modélisation physique. Ainsi, il serait

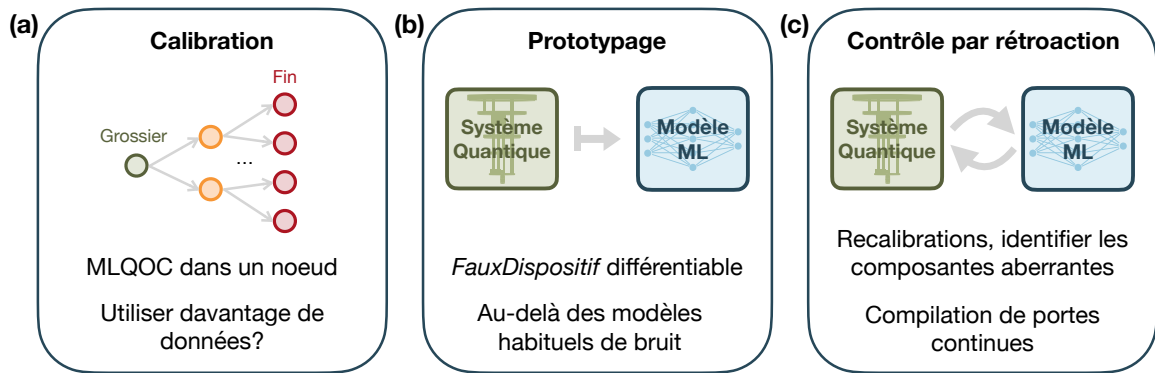


FIGURE 4.9 – Trois exemples d’applications de l’approche MLQOC. **(a)** Au-delà de l’étalonnage fin des opérations logiques, de nombreuses opérations de calibration pourraient bénéficier de l’approche MLQOC. **(b)** Le modèle entraîné représente une approximation de la transformation réalisée par le dispositif quantique. Réalisant une transformation qui est différentiable et facile à évaluer numériquement, ce modèle peut alors être utilisé dans toutes sortes de tâches de simulation telles que le prototypage de protocoles. **(c)** Le modèle entraîné peut également être utilisé en rétroaction directe avec le système quantique afin de contrôler une partie de son fonctionnement en temps réel.

intéressant de développer de nouveaux protocoles en utilisant cette représentation approximative afin de concevoir des approches qui sont robustes aux sources de bruit rencontrées dans l’expérience.

Enfin, si les conditions expérimentales le permettent, il serait également possible d’intégrer le modèle d’apprentissage dans une approche de rétroaction avec le système quantique afin de contrôler certaines opérations en temps réel. Alors que le modèle est entraîné et que ses paramètres sont fixes, une telle approche de contrôle par rétroaction permettrait par exemple d’utiliser l’information produite par le système lors de son utilisation afin de déclencher de nouveaux étalonnages suite à la dérive des paramètres du système. L’idée étant simplement qu’un modèle rapide à évaluer numériquement peut être utilisé en temps réel avec l’opération du dispositif quantique. Évidemment, il pourrait être intéressant de mettre à jour les paramètres du modèle à partir de l’information sortant du système quantique dans une telle configuration, une approche qui s’apparente aux très prometteuses méthodes d’apprentissage par renforcement basées sur un modèle [169, 150].

Prenant un pas de recul sur l’ensemble des travaux couverts dans ce chapitre, l’apprentissage automatique est un outil très puissant et utile pour résoudre toutes sortes de problèmes en informatique quantique et il est évident que ses applications vont continuer de se multiplier dans le futur. Je tiens cependant à souligner que ce ne sont pas tous les problèmes qui doivent nécessairement être résolus avec cet outil. Une compréhension heuristique du fonctionnement du système ne devrait pas remplacer la compréhension physique nécessaire à la conception de meilleurs dispositifs quantiques. Je considère plutôt que les

représentations heuristiques et fondamentalement physiques peuvent bénéficier l'une de l'autre afin de faire progresser notre compréhension et notre contrôle des technologies quantiques. Une façon de contribuer à ce co-développement serait par exemple d'étudier de nouvelles approches d'apprentissage automatique taillées sur mesure aux problèmes de nature quantique que nous rencontrons en informatique quantique.

Chapitre 5

Protocoles de caractérisation quantique

Au chapitre précédent, nous avons vu comment utiliser une approche d'apprentissage automatique afin de caractériser les dynamiques arbitraires d'un système quantique. Dans ce chapitre, nous développons des approches beaucoup plus ciblées afin de caractériser de façon précise et efficace certains éléments des dynamiques totales du système. Plus spécifiquement, nous cherchons à quantifier la qualité des opérations logiques à un et deux qubits réalisées sur des qubits supraconducteurs. Dans l'objectif d'utiliser l'information obtenue pour améliorer le fonctionnement du dispositif en pratique, nous souhaitons non seulement mesurer la fidélité des opérations logiques, mais également caractériser les sources d'erreurs expérimentales causant les imperfections mesurées.

Ce chapitre porte sur les travaux que j'ai réalisés dans le cadre de mon stage étendu avec l'équipe de Google Quantum AI. En tant qu'étudiant-chercheur au cours des deux dernières années, j'ai eu la chance de travailler directement avec les équipes théoriques et expérimentales œuvrant à la conception d'un ordinateur quantique tolérant aux fautes. Confronté aux problématiques associées à l'opération d'un dispositif supraconducteur de grande échelle, j'ai contribué à l'élaboration de protocoles permettant de caractériser les portes logiques des transmons à fréquence variable utilisés chez Google. J'ai d'ailleurs eu l'occasion de développer et de réaliser mes propres expériences sur ces qubits afin de démontrer la performance des protocoles présentés dans ce chapitre.

Dans la section qui suit, j'introduirai les problématiques ayant motivé la conception des deux protocoles de caractérisation sur lesquels j'ai principalement contribué. Ces protocoles sont intitulés *CAFE* et *MEADD*, des acronymes dont je décrirai la signification en détail.

CAFE a pour but de caractériser la fidélité d’une opération logique à un ou deux qubits et d’obtenir un bilan des contributions cohérentes et incohérentes à l’erreur totale mesurée. De son côté, MEADD a pour but de caractériser les paramètres unitaires décrivant ces mêmes opérations. Je présenterai les résultats principaux associés à ces deux protocoles à la section 5.2, avant de reproduire les articles scientifiques les introduisant à la communauté aux sections 5.3 et 5.4. Je conclurai avec une discussion ouverte de différentes applications et extensions de ces protocoles à la section 5.5.

5.1 Mise en contexte

L’élaboration et l’opération d’un ordinateur quantique à grande échelle, possédant aujourd’hui une centaine de qubits, mais éventuellement des millions de qubits, font naître de nombreuses problématiques qui ne sont typiquement pas rencontrées dans le contexte d’expérimentations sur quelques qubits. Dans ma recherche, je m’intéresse principalement à comprendre le fonctionnement des qubits supraconducteurs et à utiliser cette information afin d’améliorer les opérations réalisées en pratique. En appliquant ces tâches de caractérisation et de contrôle à un dispositif possédant de nombreux qubits, les principales problématiques proviennent sans surprise de la complexité des dynamiques quantiques dans un grand espace d’Hilbert, mais également de la stabilité des opérations dans le temps. En effet, les conditions expérimentales ont tendance à changer de façon significative sur une période d’environ 24 h pour diverses raisons plus ou moins bien comprises telles que des changements en température des appareils électroniques de contrôle [170] et des mouvements en fréquence des défauts à deux niveaux présents dans les matériaux entourant les qubits [82]. Ainsi, la caractérisation et la calibration des opérations sont typiquement répétées à tous les jours sur l’ensemble des qubits, ce qui pose des contraintes importantes sur la rapidité d’exécution et d’analyse des protocoles utilisés, ainsi que sur leur robustesse devant des qubits et des opérations ayant des propriétés et des sources de bruit pouvant varier considérablement.

Suite à ces contraintes d’efficacité et de robustesse des protocoles assurant leur bon fonctionnement à grande échelle, on comprend qu’il devient essentiel de concevoir des approches de caractérisation spécialisées afin d’extraire l’information utile au contrôle en utilisant le moins de ressources possible. Une idée simple mais importante est que certaines informations sur la performance du dispositif sont plus utiles que d’autres dans le but d’améliorer cette même performance en pratique. Par exemple, l’étalonnage de la performance globale du système avec une mesure comme le volume quantique [171], l’entropie croisée associée à l’échantillonnage de circuits aléatoires [172], ou le taux d’erreur logique

d'un qubit encodé dans un code de surface [1, 15] sont toutes des métriques très utiles pour comparer différents systèmes entre eux, mais s'avèrent peu utiles pour indiquer les actions à entreprendre afin d'améliorer cette même métrique. Dans ce qui suit, nous introduisons deux protocoles, CAFE et MEADD, spécifiquement conçus pour répondre aux besoins en efficacité, en robustesse et en informations pertinentes pour le contrôle.

5.1.1 Motivations pour CAFE

Une approche communément utilisée pour obtenir de l'information pertinente à l'amélioration du contrôle consiste à caractériser individuellement la fidélité de chacune des opérations utilisées. Cette méthode permet de diagnostiquer les opérations les plus fautives pour ensuite les calibrer à nouveau. Une limitation importante provenant d'une telle approche est qu'en raison des imperfections, le contexte spatial et temporel dans lequel l'opération est réalisée peut grandement affecter sa fidélité. Ainsi, la performance de l'opération dans le contexte utilisé lors de sa caractérisation n'est pas nécessairement représentative de sa contribution à l'erreur totale du circuit d'intérêt. Cette idée de l'importance du contexte spatial et temporel sur la fidélité d'une opération est illustrée à la Fig. 5.1.

Par contexte spatial, je réfère ici à l'environnement physique dans lequel le ou les qubits se trouvent lors de l'opération d'intérêt. Comme discuté au chapitre 4, la diaphonie classique (provenant des contrôles appliqués sur d'autres qubits) et quantique (provenant de l'échange d'information entre les qubits voisins en raison d'un couplage non nul) fait en sorte que la fidélité d'une opération en isolation peut être très différente de celle réalisée en parallèle avec d'autres opérations. Par exemple, l'erreur moyenne d'une porte CZ réalisée en parallèle avec d'autres CZ est typiquement deux fois plus grande que lorsque la même opération est exécutée en isolation [172]. De façon similaire, l'importance du contexte temporel des opérations provient du fait qu'en pratique l'environnement des qubits possède une mémoire finie, ce qui peut faire en sorte que l'action précise d'un signal de contrôle devienne dépendante de l'historique des signaux appliqués précédemment. Pour les qubits supraconducteurs, cette dépendance temporelle provient notamment des signaux parasites dans les lignes de contrôle [70, 97] et des excitations résiduelles dans la cavité micro-onde suite à une mesure du qubit [173].

Somme toute, on comprend qu'il est désirable de caractériser les opérations en conservant le plus possible le contexte dans lequel elles sont exécutées dans le circuit quantique d'intérêt. C'est précisément ce qu'on cherche à faire avec notre protocole intitulé *Context-aware fidelity estimation* (CAFE) [3], lequel permet d'obtenir la fidélité d'une opération ou d'un groupe d'opérations quantiques. En plus de cette idée de caractériser les opérations dans le même

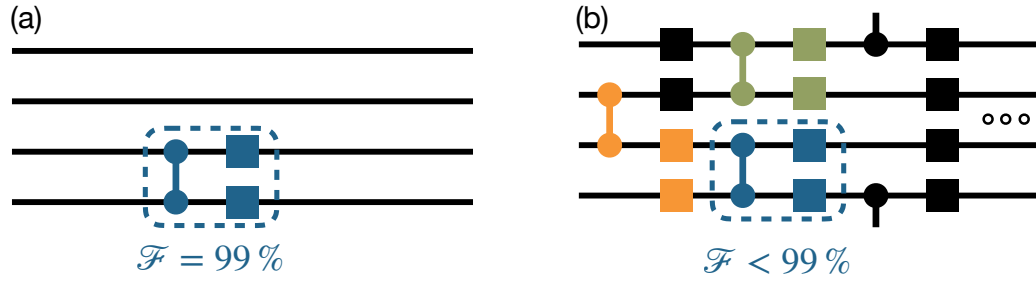


FIGURE 5.1 – La fidélité d’une opération quantique peut dépendre significativement du contexte spatial et temporel dans lequel elle se trouve. (a) Un groupe d’opérations (bleu) réalise la transformation unitaire souhaitée avec une certaine fidélité moyenne $\mathcal{F}$. (b) La fidélité du même groupe d’opération (bleu) peut être significativement différente lorsqu’il est introduit dans le circuit quantique d’intérêt comportant d’autres opérations. Notamment, le contexte temporel défini par les opérations précédentes (orange) et le contexte spatial défini par les opérations réalisées en parallèle (vert) peuvent tous deux interférer avec le groupe d’opérations de sorte à réduire sa fidélité.

contexte où on veut les utiliser, CAFE est conçu pour ne nécessiter que très peu d’expériences, pour être robuste aux erreurs associées à la préparation et à la mesure des qubits, ainsi que pour fournir un budget de l’erreur totale en termes de contributions cohérentes (provenant de processus unitaires sur système d’intérêt) et incohérentes (provenant de pertes stochastiques associées à l’échange d’information avec les degrés de liberté de l’environnement).

Cette distinction entre la présence de processus parasites qui sont cohérents ou incohérents est importante alors que les actions à entreprendre en pratique pour les atténuer peut grandement différer. En effet, une source d’erreur cohérente peut être décrite par une transformation unitaire (laquelle correspond à une rotation de l’état du qubit ou du système quantique d’un angle et d’un axe donné) et peut donc être corrigée directement en utilisant notre contrôle universel de l’état du système. En pratique, cela correspond typiquement à ajuster la valeur des paramètres définissant la forme des signaux de contrôle utilisé de sorte à minimiser cette erreur. D’un autre côté, les erreurs incohérentes ne peuvent pas de façon générale être éliminées avec des contrôles cohérents sur le système quantique. Cependant, il est possible d’atténuer l’amplitude des erreurs incohérentes à partir de contrôles cohérents sur le système. L’exemple le plus simple d’une telle approche pour les qubits supraconducteurs à fréquence variable consiste à ajuster la fréquence d’opération du ou des transmons afin d’exploiter un régime où les pertes sont réduites. Le découplage dynamique [174], qui sera présenté plus en détail un peu plus bas, est un autre exemple de l’utilisation de contrôles cohérents pour atténuer l’impact de processus incohérents.

En utilisant les approches les plus répandues pour caractériser la fidélité d’opérations quantiques, telles que l’étalonnage par portes aléatoires (*randomized benchmarking*, RB) [175, 45] et la mesure d’entropie croisée (*cross entropy benchmarking*, XEB) [176, 172], il s’avère

cependant ardu de quantifier précisément les contributions cohérentes et incohérentes à l'erreur moyenne totale. Cette difficulté provient du fait que la grande majorité des protocoles de caractérisation introduisent des opérations aléatoires et moyennent sur les résultats afin d'estimer la fidélité d'un ensemble cible d'opérations (ou d'une seule opération) [177]. De façon effective, cette moyenne permet d'uniformiser les différentes sources de bruits de sorte que l'erreur totale puisse être amplifiée en utilisant des séquences de portes plus longues puis approchée par une simple exponentielle. Par exemple, en échantillonnant les opérations aléatoires à partir du groupe de Clifford, l'action combinée de l'opération d'intérêt et des portes aléatoires peut être décrite par un canal dépolarisant [44]. L'effet d'un tel canal est simplement de rétrécir la sphère de Bloch de façon uniforme [62]. Ainsi, la fidélité de l'opération d'intérêt peut être estimée en mesurant la probabilité de réaliser l'opération voulue en fonction du nombre d'applications de cette opération m , puis en ajustant ces probabilités avec une fonction exponentielle de la forme

$$p(m) = Af^m + \frac{1}{2^n}, \quad (5.1)$$

pour une opération sur n qubits où $\{A, f\}$ sont des paramètres d'ajustement tels que A représente les erreurs de préparation et de mesure (SPAM) et f est l'estimation de la fidélité moyenne de l'opération. Cette idée d'ajouter des portes aléatoires afin de simplifier la forme effective du bruit moyen est au cœur de nombreux protocoles de caractérisation quantique [177, 162].

Par le fait même d'introduire des portes aléatoires, l'information sur les dynamiques cohérentes est cependant masquée dans l'erreur totale rendue incohérente. Il devient donc difficile d'estimer de façon précise sa contribution. Quelques approches existent pour y arriver, lesquelles utilisent principalement une estimation indépendante de l'erreur incohérente puis combinent ces estimés bruités afin d'obtenir la contribution unitaire à l'erreur totale telle que $e_{\text{coh.}} = e_{\text{tot.}} - e_{\text{incoh.}}$. Cette estimation de l'erreur incohérente se fait typiquement en mesurant la pureté de l'état quantique dans les expériences de RB [178] ou en analysant la variance des distributions en XEB [172]. Cependant, cette simple combinaison des erreurs incohérentes et totales mesurées fait en sorte que l'estimation de l'erreur cohérente possède une grande incertitude et fournit fréquemment des résultats non physiques avec des valeurs négatives.

Avec CAFE, nous n'introduisons aucune porte aléatoire afin de préserver la structure originale du bruit associé au groupe d'opérations à caractériser. Cela nous permet d'amplifier directement les erreurs cohérentes et incohérentes de sorte à pouvoir estimer leurs contributions de manière harmonieuse. En effet, les erreurs incohérentes $e_{\text{incoh.}}$ proviennent de

processus stochastiques et s'accumulent donc par définition de façon linéaire avec le nombre d'exécutions de l'opération m . De l'autre côté, les erreurs cohérentes $e_{\text{coh.}}$ sont des processus unitaires agissant sur l'état quantique d'intérêt. Comme la fidélité de porte moyenne est une fonction quadratique de l'état quantique évolué suite à l'opération, voir l'Éq. (2.25), les erreurs cohérentes s'accumulent plutôt de façon quadratique avec m au premier ordre. Comme nous le verrons dans l'article introduisant CAFE plus en détail, il est alors possible d'utiliser cette différente dépendance en m afin d'estimer ces deux quantités, $e_{\text{incoh.}}$ et $e_{\text{coh.}}$, de manière cohérente directement à partir des données fournies par notre protocole.

5.1.2 Motivations pour MEADD

Une fois qu'une méthode comme CAFE a été utilisée afin de diagnostiquer les opérations majoritairement responsables de l'erreur totale observée, il est nécessaire d'ajuster nos contrôles afin d'améliorer la performance de ces opérations en pratique. Pour ce faire, il est particulièrement utile de connaître avec précision quel paramètre unitaire de l'opération cible doit être ajusté. Une telle caractérisation permet d'éviter d'avoir à balayer chacun des paramètres de contrôle de sorte à améliorer la fidélité mesurée, un processus de calibration nécessitant de nombreuses expériences [116, 117].

L'une des opérations limitant principalement la performance des processeurs supraconducteurs s'avère être les portes logiques à deux qubits [1, 15]. Pour les dispositifs conçus à Google, ces opérations sont décrites par un ensemble de portes conservant le nombre total d'excitations dans les deux qubits. Cet ensemble est communément appelé *fSim* pour simulation fermionique, alors que la conservation du nombre de photons correspond directement à la conservation du nombre d'électrons dans la simulation de systèmes quantiques tels que ceux utilisés en chimie quantique [101]. Cet ensemble de portes est réalisé en contrôlant le flux externe traversant la boucle supraconductrice des transmons et de leur coupleur [13]. Ignorant les phases associées à chacun des deux qubits, l'unitaire de ces opérations s'exprime comme [101]

$$\text{fSim}(\theta, \phi) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & \cos \theta & -i \sin \theta & 0 \\ 0 & -i \sin \theta & \cos \theta & 0 \\ 0 & 0 & 0 & e^{-i\phi} \end{pmatrix}, \quad (5.2)$$

où θ est l'angle associé à l'échange d'une excitation entre les deux qubits $|01\rangle \leftrightarrow |10\rangle$ et ϕ

est la phase conditionnelle accumulée par l'état $|11\rangle$. L'opération principalement utilisée en correction d'erreur quantique est $\text{CZ} = \text{fSim}(0, \pi)$, bien que d'autres constructions existent basées par exemple sur la porte $\text{iSWAP} = \text{fSim}(\pi/2, 0)$ [179]. Je note que l'opération CNOT communément utilisée en informatique quantique est équivalente au CZ avec l'ajout, sur le qubit cible, d'une porte Hadamard précédent et suivant l'opération à deux qubits.

Alors que les phases à un qubit peuvent être facilement contrôlées en ajustant la phase des signaux micro-ondes utilisés lors des opérations à un qubit précédant et suivant la porte fSim [180], les angles θ et ϕ définissent l'intrication entre les qubits et dépendent donc des dynamiques lors de l'opération à deux qubits. Afin d'être en mesure de calibrer des opérations à deux qubits de haute fidélité ($> 99.9\%$), il est alors souhaitable de caractériser ces deux paramètres avec précision. En effet, une erreur de seulement 3 degrés (0.05 rad) sur ces angles mène à une fidélité de porte moyenne inférieure à 99.85 % en l'absence de toute autre erreur.

L'approche communément utilisée afin d'estimer la valeur des paramètres unitaires d'une opération consiste à répéter cette opération un nombre variable de fois afin d'amplifier les erreurs et ainsi obtenir un signal mesurable. Comme les angles de rotations et les phases définis comme paramètres unitaires s'accumulent de façon cohérente avec chaque répétition, il est possible d'ajuster les données obtenues avec un modèle simple et d'obtenir le paramètre associé à une seule réalisation de l'opération cible avec une bonne précision. Par exemple, pour estimer l'angle de rotation d'une porte à un qubit $R_x(\pi)$ pour une opération imparfaite réalisée dans l'expérience $\hat{U}_{\text{exp.}} = R_x(\pi + \lambda)$, il suffit de l'appliquer un nombre pair de fois $2m$ à l'état initial $|0\rangle$ pour différents $m \in \mathbb{N}$ puis de mesurer le qubit dans la base de calcul afin d'obtenir la probabilité

$$\text{Prob}(m) = |\langle 1 | (\hat{U}_{\text{exp.}})^{2m} | 0 \rangle|^2 \quad (5.3)$$

$$= |\langle 1 | e^{-im(\pi+\lambda)\hat{\sigma}_x} | 0 \rangle|^2 \quad (5.4)$$

$$= |\langle 1 | e^{-im\lambda\hat{\sigma}_x} e^{-im\pi\hat{\sigma}_x} | 0 \rangle|^2 \quad (5.5)$$

$$= |(-1)^m \langle 1 | e^{-im\lambda\hat{\sigma}_x} | 0 \rangle|^2 \quad (5.6)$$

$$= |\langle 1 | \cos(m\lambda)\mathbb{1} - i\sin(m\lambda)\hat{\sigma}_x | 0 \rangle|^2 \quad (5.7)$$

$$= \sin^2(m\lambda) \quad (5.8)$$

Cette expérience produit alors des oscillations de Rabi qui peuvent aisément être ajustées à un modèle d'oscillations sinusoïdales afin d'obtenir le paramètre λ recherché. Ce type d'approche par amplification de l'amplitude peut être utilisée de façon générale afin d'esti-

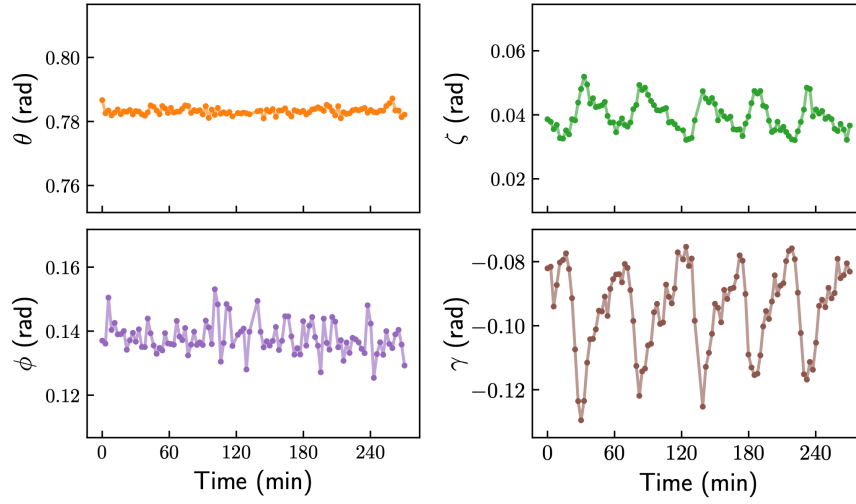


FIGURE 5.2 – Fluctuations temporelles des paramètres unitaires d’une opération CZ telles que caractérisées par une approche de Floquet [170]. Les paramètres ζ et γ associés à la phase différentielle et commune des qubits, lesquelles peuvent toutes deux être contrôlées par des opérations à un qubit, varient de façon significative avec le temps et semblent osciller périodiquement. Bien que mal compris, ce phénomène est typiquement attribué à des fluctuations en température des appareils électroniques de contrôle et/ou de l’environnement des qubits. Figure tirée de la Réf. [170].

mer les paramètres unitaires décrivant la porte logique d’intérêt. Il suffit alors d’adapter les états quantiques préparés, les différents nombres de répétitions utilisés, la base de mesure ainsi que l’analyse subséquente des données. Par exemple, les protocoles nécessaires pour caractériser les cinq paramètres unitaires de portes logiques à deux qubits préservant le nombre total d’excitations sont décrits dans les annexes de la Réf. [170]. Ces protocoles ont d’ailleurs été nommés "calibration par Floquet", où la référence à la théorie de Floquet [181] provient de la périodicité des circuits quantiques d’amplification utilisés.

Cette approche d’amplification possède cependant des limitations importantes provenant du fait que la répétition des opérations n’amplifie pas que le paramètre unitaire à estimer, mais amplifie également le bruit ainsi que d’autres processus unitaires dont la stabilité temporelle moindre peut affecter la précision du protocole de caractérisation. Pour les qubits supraconducteurs à fréquence accordable par le flux magnétique, tels que ceux utilisés à Google, le bruit sur la fréquence de transition des qubits est une source importante d’erreurs. En raison de ce degré de liberté fortement couplé aux variations de l’environnement du qubit, la phase des qubits varie significativement dans le temps. Cet effet est illustré à l’aide de données expérimentales à la Fig. 5.2.

En plus de créer du déphasage venant réduire la visibilité du signal amplifié lors de l’utilisation de longues séquences de portes, ces variations temporelles importantes de

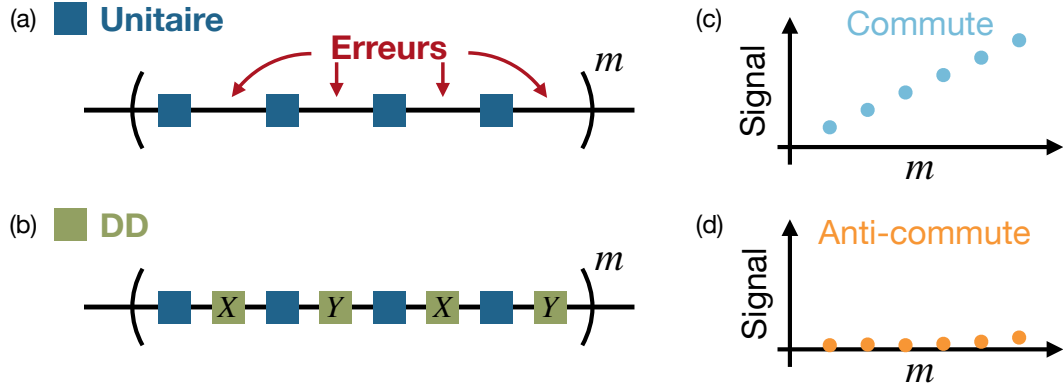


FIGURE 5.3 – Utilisation de découplage dynamique (DD) pour l'estimation de paramètres unitaires. (a) La répétition de l'unitaire cible permet d'amplifier le paramètre unitaire à estimer, mais amplifie également les erreurs parasites. (b) L'ajout de portes de découplage dynamique réalisant l'identité, par exemple la séquence $XY^4 \equiv XYXY$ [182], permet de filtrer une portion de ces erreurs. (c) Les erreurs et autres processus unitaires qui commutent avec les portes de découplage dynamique sont amplifiés de façon cohérente, alors que (d) ceux qui anti-commutent sont éliminés au premier ordre.

la phase des qubits peuvent directement brouiller le signal moyen mesuré. C'est le cas par exemple de la phase conditionnelle ϕ , qui je le rappelle est l'un des deux paramètres d'intrication d'une porte fSim et qui est donc en principe totalement indépendant des phases individuelles des qubits. À la Fig. 5.2, on observe que l'estimation de ce paramètre est beaucoup moins précise que celle de θ , l'autre paramètre d'intrication. Ceci est dû au fait que la phase totale accumulée par l'état $|11\rangle$ lors d'une porte fSim est donnée par $\phi + 2\gamma$. Ainsi, le protocole de Floquet produit un signal venant amplifier à la fois le paramètre stable ϕ que l'on veut estimer ainsi que la phase à un qubit γ qui fluctue significativement dans le temps. De ce fait, la capacité du protocole à estimer ϕ de façon précise est fortement limitée.

Afin de surmonter ces limitations de l'estimation de paramètres unitaires par amplification, nous introduisons un protocole faisant appel au découplage dynamique. Cette approche de caractérisation, que nous intitulons MEADD pour *Matrix Element Amplification using Dynamical Decoupling* [4], sera décrite en détail à la section 5.4. La principale intuition pour bien comprendre en quoi le découplage dynamique est particulièrement bénéfique à notre approche de caractérisation est la suivante : l'insertion d'une séquence de portes quantiques à un qubit se combinant à l'identité dans un circuit répété permet d'éliminer au premier ordre tous les processus unitaires qui anti-commutent avec ces portes. Cette idée est illustrée de façon schématique à la Fig. 5.3. Avec MEADD, les processus unitaires que nous éliminons sont tous ceux pouvant interférer avec le paramètre unique que l'on souhaite estimer.

Afin de bien comprendre cette idée, étudions un exemple simple s'apparentant à une

expérience d'écho de spin. Ici, on considère la situation artificielle où l'on souhaite caractériser une faible rotation ϵ_x autour de l'axe X d'une opération. Supposons que cette opération expérimentale réalise plutôt des rotations parasites à la fois autour de X et de Y de sorte à être décrite par

$$\hat{U}_{\text{exp.}} = R_x(\epsilon_x)R_y(\epsilon_y) = e^{-i\epsilon_x\hat{\sigma}_x/2}e^{-i\epsilon_y\hat{\sigma}_y/2}. \quad (5.9)$$

Avec une approche directe d'amplification telle que celle décrite plus haut, le signal produit dépend à la fois de ϵ_x et de ϵ_y alors que

$$\text{Prob}(m) = |\langle 1 | (\hat{U}_{\text{exp.}})^m | 0 \rangle|^2 \quad (5.10)$$

$$= \left| \langle 1 | e^{-im\epsilon_x\hat{\sigma}_x/2} e^{-im\epsilon_y\hat{\sigma}_y/2} + \mathcal{O}(\epsilon_j^2) | 0 \rangle \right|^2 \quad (5.11)$$

$$\approx m^2(\epsilon_x + \epsilon_y)^2/4, \quad (5.12)$$

pour $\epsilon_x, \epsilon_y \ll 1$. On comprend donc que notre estimation du paramètre ϵ_x va être biaisée par la rotation parasite ϵ_y . En ajoutant maintenant un nombre pair de portes de découplage dynamique $\hat{\sigma}_x$ de sorte que l'opération cycle répétée m fois devienne

$$\hat{U}_{\text{cycle}} = \hat{U}_{\text{exp.}}\hat{\sigma}_x\hat{U}_{\text{exp.}}\hat{\sigma}_x, \quad (5.13)$$

nous sommes en mesure d'éliminer la dépendance en ϵ_y alors que l'opérateur $\hat{\sigma}_y$ anti-commute avec les portes de découplage dynamique $\hat{\sigma}_x$. En effet,

$$\hat{U}_{\text{exp.}}\hat{\sigma}_x = \hat{\sigma}_x e^{-i\epsilon_x\hat{\sigma}_x/2} e^{i\epsilon_y\hat{\sigma}_y/2}, \quad (5.14)$$

de sorte que

$$\hat{U}_{\text{cycle}} = e^{-i\epsilon_x\hat{\sigma}_x} + \mathcal{O}(\epsilon_j^2). \quad (5.15)$$

On retrouve donc directement le signal $\text{Prob}(m) = \sin^2(m\epsilon_x)$ voulu grâce au découplage dynamique. Avec MEADD, on applique cette même idée afin d'estimer les paramètres des portes fSim avec une précision surpassant tout autre approche.

5.2 Résultats principaux

Dans les articles scientifiques qui suivent, nous introduisons deux protocoles de caractérisation permettant principalement d’estimer la fidélité de portes logiques sur deux qubits et d’obtenir leur description unitaire complète. À l’aide de simulations et de réalisations expérimentales, nous démontrons que ces protocoles surpassent les approches existantes en termes d’efficacité expérimentale et de précision. Ces deux protocoles de caractérisation sont maintenant utilisés quotidiennement dans la calibration des opérations sur les dispositifs à Google Quantum AI et peuvent être utilisées directement ou facilement adaptées afin d’être appliqués à de nombreux autres dispositifs quantiques.

CAFE introduit une approche expérimentale minimale mais complète afin de caractériser une opération ou d’un groupe d’opérations dans le contexte même de leur circuit quantique. Ce protocole permet d’estimer précisément la fidélité de porte moyenne ainsi que les contributions cohérentes (unitaires) et incohérentes (stochastiques) à l’erreur totale mesurée. Nos simulations de portes CZ avec des taux d’erreurs réalistes démontrent que CAFE est légèrement plus précis que l’approche la plus répandue (RB), tout en utilisant significativement moins de ressources expérimentales. En utilisant différentes répétitions du circuit d’intérêt et en ajustant les fidélités moyennes mesurées à un modèle, CAFE est robuste aux erreurs de préparation d’état et de mesure (SPAM).

Nous appliquons CAFE expérimentalement afin de démontrer l’évaluation précise des erreurs cohérentes et incohérentes d’une centaine de portes CZ sur un dispositif Sycamore [172]. Nous utilisons également CAFE afin d’estimer la fidélité de plusieurs groupes d’opérations dans le contexte du circuit utilisé pour stabiliser l’état quantique d’un code de surface de distance trois. Dans cette situation, CAFE nous permet d’observer une signature claire de certaines erreurs présentes uniquement lors de la réalisation concurrente des autres portes quantiques, illustrant ainsi l’utilité de caractériser les opérations dans le contexte où elles se trouvent.

Enfin, CAFE peut directement être utilisé afin de valider notre compréhension des erreurs cohérentes dans les opérations réalisées. En effet, il est possible d’effectuer les mêmes expériences, mais de changer les portes logiques utilisées à la toute fin avant de mesurer l’état quantique final afin de tester la fidélité de l’opération réalisée par rapport à une unitaire arbitraire. Par exemple, il est possible de mesurer la fidélité d’une réalisation donnée par rapport à l’unitaire cible CZ, mais également de mesurer la fidélité par rapport à une autre unitaire comme $\text{fSim}(0.01, \pi - 0.02)$. À ce moment, il est possible de comparer la contribution cohérente à l’erreur totale entre ces deux caractérisations afin de vérifier si l’unitaire

$\text{fSim}(0.01, \pi - 0.02)$ décrit bel et bien l'opération expérimentale mieux que l'unitaire cible $\text{CZ} = \text{fSim}(0, \pi)$. Nous utilisons directement cette propriété du protocole CAFE afin de démontrer que la caractérisation unitaire obtenue par MEADD est significativement plus exacte que toute autre approche existante.

Avec notre second protocole, MEADD, nous surmontons la principale limitation à la caractérisation précise des paramètres unitaires des portes logiques réalisées sur des qubits supraconducteurs à fréquence variable : le bruit en fréquence. Pour ce faire, nous utilisons une approche de découplage dynamique (DD) où nous ajoutons des portes à un qubit entre chacune des répétitions de l'opération cible. Au-delà de l'intuition conventionnelle selon laquelle le découplage dynamique filtre le bruit de basse fréquence [183], nous introduisons l'idée d'utiliser le DD afin de filtrer spécifiquement certains éléments de matrice d'une transformation unitaire. À partir d'une conception soignée des circuits quantiques utilisés dans le protocole, il est alors possible d'amplifier de façon cohérente le ou les paramètres que l'on cherche à estimer tout en éliminant l'impact des autres termes ainsi que des sources de bruit de basse fréquence. De manière cruciale, ces estimations sont robustes au premier ordre à toutes les erreurs associées aux portes à un qubit que l'on ajoute afin de filtrer l'information que l'on souhaite amplifier.

En nous concentrant sur les portes CZ, nous démontrons expérimentalement sur un processeur Sycamore que MEADD est cinq fois plus précis que les approches existantes pour estimer la phase conditionnelle ϕ et dix fois plus précis pour estimer l'angle d'échange d'excitation θ . Pour l'angle θ , cela correspond à une incertitude de moins de 0.4 mrad (0.02°), une réalisation faisant de MEADD un outil indispensable dans la calibration d'opérations quantiques à deux qubits de très haute fidélité.

Nous démontrons enfin l'utilité de MEADD dans la caractérisation d'opérations à un qubit ainsi que l'estimation de la diaphonie quantique entre deux qubits voisins impliqués dans des portes CZ parallèles. Pour les portes à un qubit, MEADD est plus précis que l'approche répandue par un facteur de plus de six. Concernant la diaphonie quantique, MEADD est en mesure d'estimer en médiane des rotations de (0.42 ± 0.07) mrad alors que ces processus parasites sont sous le seuil de bruit des autres approches et donc indétectables en pratique. Nous fournissons également la formulation de protocoles MEADD afin de caractériser d'autres portes à deux qubits préservant le nombre total d'excitations telles que iSWAP et $\sqrt{\text{iSWAP}}$.

5.3 Publication [3] : Caractérisation de la fidélité des opérations quantiques (CAFE)

Je suis premier auteur à contribution égale avec deux autres collaborateurs de l'article suivant publié dans la revue *Physical Review Research*. Cette attribution des contributions témoigne du fort environnement de collaboration dans lequel ces travaux ont été réalisés. En effet, j'ai activement contribué à la détermination analytique du modèle d'ajustement à utiliser pour caractériser les contributions cohérentes et incohérentes des portes à un des deux qubits, à la mise en œuvre numérique du protocole ainsi qu'à la rédaction du manuscrit. J'ai réalisé l'ensemble des simulations numériques, en plus d'acquérir les données expérimentales, d'analyser les résultats et de produire toutes les figures.

Je note qu'en raison d'un problème technique causé par quelques figures trop volumineuses numériquement en format PDF, l'article qui suit n'est pas la version publiée dans PRR, laquelle est disponible à la Réf. [3], mais une reproduction exacte où la résolution des figures a été réduite.

Context-Aware Fidelity Estimation

Dripto M. Debroy,^{1,*} Élie Genois,^{1,2,†} Jonathan A. Gross,^{1,‡} Wojciech Mruczkiewicz,¹
Kenny Lee,¹ Sabrina Hong,¹ Zijun Chen,¹ Vadim Smelyanskiy,¹ and Zhang Jiang¹

¹Google Quantum AI, Venice, CA 90291

²Institut quantique & Département de Physique, Université de Sherbrooke, Québec J1K 2R1, Canada

We present Context Aware Fidelity Estimation (CAFE), a framework for benchmarking quantum operations that offers several practical advantages over existing methods such as Randomized Benchmarking (RB) and Cross-Entropy Benchmarking (XEB). In CAFE, a gate or a subcircuit from some target experiment is repeated n times before being measured. By using a subcircuit, we account for effects from spatial and temporal circuit context. Since coherent errors accumulate quadratically while incoherent errors grow linearly, we can separate them by fitting the measured fidelity as a function of n . One can additionally interleave the subcircuit with dynamical decoupling sequences to remove certain coherent error sources from the characterization when desired. We have used CAFE to experimentally validate our single- and two-qubit unitary characterizations by measuring fidelity against estimated unitaries. In numerical simulations, we find CAFE produces fidelity estimates at least as accurate as Interleaved RB while using significantly fewer resources. We also introduce a compact formulation for preparing an arbitrary two-qubit state with a single entangling operation, and use it to present a concrete example using CAFE to study CZ gates in parallel on a Sycamore processor.

I. INTRODUCTION

Reliably understanding the structure of noise in quantum processors is vital for advancing quantum computation. There are a variety of characterization methods available to quantify and validate the performance of individual quantum operations like state preparation, single- and two-qubit gates, measurement, and reset [1]. Amongst these techniques, some use randomness to estimate gate fidelities efficiently, such as randomized benchmarking (RB) [2–7], cross-entropy benchmarking (XEB) [8, 9], channel spectrum benchmarking [10], and direct fidelity estimation [11–13], while others use complete sets of input states, such as unitary tomography [14, 15], quantum process tomography [16, 17], and gate set tomography [18, 19]. These methods all have upsides and downsides in terms of speed, scalability, and the amount of information provided [1].

An important problem that many existing techniques face is that as quantum computers scale, complex inter-component interactions arise, such as control crosstalk or temporal correlations caused by residual pulse tails. For example, experiments may use amplifiers with temperature dependent gain profiles, which could be linked to pulse duty cycle and cause circuit-dependent errors. These effects can make the performance of a quantum operation highly context dependent [9, 20–22], and there is a growing need for characterization methods which account for them.

In this paper, we describe a characterization method called *Context-Aware Fidelity Estimation* (CAFE) which

measures the fidelity between an experimentally implemented quantum operation and a reference unitary. In CAFE, we repeat a gate or a subcircuit from a target experiment n times and measure the average gate fidelity against the reference unitary raised to the n^{th} power. For example, the subcircuit can be part of a stabilizer extraction circuit in a quantum error correction experiment, as demonstrated experimentally in Sec. VI. This procedure allows us to capture context-related errors which isolated characterizations do not capture [23, 24]. Since coherent errors accumulate quadratically as a function of n while incoherent errors grow linearly, we can separate their contributions to the infidelity. This is in contrast with RB and XEB, where random compiling is used to twirl coherent errors into incoherent ones.

Throughout the main text of this paper, we focus on using CAFE to study the performance of CZ gates implemented in parallel, as two-qubit operations have been shown to be a dominant error source across recent large-scale experiments on a number of experimental platforms, especially in the presence of stray interactions [25–27]. We also discuss CAFE for a subcircuit of a stabilizer extraction circuit. In the Appendices, we present results from CAFE experiments characterizing parallel single-qubit operations.

II. CONTEXT AWARE FIDELITY ESTIMATION

A typical CAFE experiment is performed in three steps, schematically shown in Fig. 1. First, a state $|\psi\rangle$ is prepared from an m -qubit 2-design $\{\psi_i\}$ [28–31]. These ensembles match the Haar-random distribution up to second moments, allowing us to measure fidelity. In App. B, we present a method to construct shallow circuits that prepare arbitrary two-qubit entangled states for this purpose. The method uses a single entangling operation to

* dripto@google.com; Contributed Equally

† elie.genois@usherbrooke.ca; Contributed Equally

‡ jarthurgross@google.com; Contributed Equally

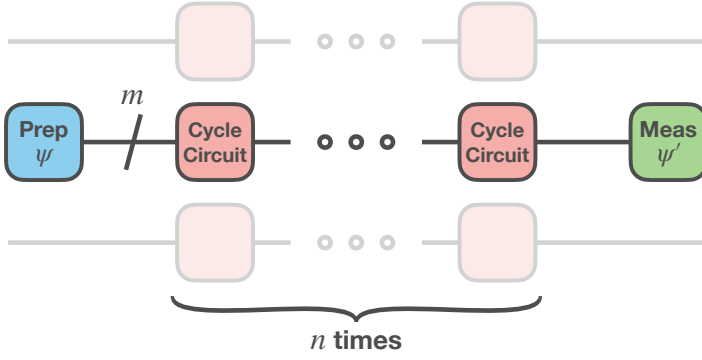


FIG. 1. A schematic of the circuits used to measure the fidelity of n repetitions of the cycle circuit, which is selected to be a subcircuit from some target experiment. First, a state pulled from a m qubit 2-design, $\{\psi_i\}$, is prepared (blue). Second, the cycle circuit being characterized is applied n times (pink). Third, we undo the combined action of the previous circuit in a single step, similar to the inversion step in randomized benchmarking, and measure the resulting state in the computational basis (green). The fidelity between n repetitions of the applied operation and the n^{th} power of the reference unitary can be found by averaging the probability of getting $|0\rangle^{\otimes m}$ over all 4^m initial states. The faded operations represent the spatial context of the cycle circuit being characterized.

prepare a two-qubit state with the desired entanglement signature, followed by single-qubit operations to reach the desired state. The structure of the circuits produced by this method is illustrated in Fig. 3(a) in blue.

Next, the m -qubit circuit of interest, referred to as the *cycle circuit*, is repeated n times, along with operations on neighboring uninvolved qubits. The cycle circuit can include multiple operations, allowing close matching to the circuit context found in the target experiment. As an example, an experiment which includes dynamical decoupling (DD) would be robust to certain coherent errors, and including these DD gates in the cycle circuit allows CAFE to neglect contributions from these coherent errors, focusing on the errors which impact performance. Inserting DD gates is one of several ways for CAFE to separate coherent and incoherent contributions to gate errors, which we discuss further in Sec. IV.

Lastly, we apply a circuit that maps the state back to $|0\rangle^{\otimes m}$, assuming the cycle circuit implements the desired “fiducial unitary”, before measuring the qubits in the computational basis. This final step is analogous to the inversion step of RB, used to undo the action of the total applied circuit before measurement. Moreover in CAFE, this final step allows us to validate the performance of different unitary characterization methods, as discussed further in Sec. V. The average gate fidelity of the operation, which we refer to as *fidelity* in the following, can be found by averaging over the experiments for all 2-design states $\{\psi_i\}$ [14]:

$$\mathcal{F}(U_{\text{cycle}}, U_{\text{fiducial}}) = \langle \mathbb{P}_{0\dots 0} \rangle_{\{\psi_i\}}. \quad (1)$$

We note that in the $m = 2$ case, the preparation and measurement stages only use a single entangling gate, so their contribution to the total error rate is far smaller than the circuit being characterized for most values of n . In general, since the information we extract from CAFE comes from fitting fidelity over different values of n , reason-

$$\begin{aligned} \psi &\in \{\psi_i\} \\ |\psi'\rangle &= (U_{\text{fiducial}})^n |\psi\rangle \\ \mathcal{F}_n(U_{\text{cycle}}, U_{\text{fiducial}}) &= \langle \mathbb{P}_{0\dots 0} \rangle_{\{\psi_i\}} \end{aligned}$$

able state-preparation and measurement (SPAM) errors do not significantly impact the characterization results, as they only cause a constant offset in the fidelity curve. While the number of 2-design states scales exponentially with qubit count, random circuits can approximate these states in polynomial depth [29], which could be used for performing CAFE on $m > 2$ qubits.

III. COMPARING CAFE WITH RANDOMIZED BENCHMARKING

In this section, we compare our CAFE approach to the widely used Interleaved Randomized Benchmarking (IRB) protocol for estimating the fidelity of noisy CZ gates. The error model for these gates, which is fully described in App. E, contains coherent errors together with amplitude and phase damping as incoherent noise. In Fig. 2, we show that for this model, and using realistic error rates, CAFE yields a more accurate average gate fidelity estimation than IRB, while requiring significantly less experimental resources.

More concretely, the randomized benchmarking protocol uses 20 different circuits for 7 depths $5 \leq n \leq 35$ of random two-qubit Cliffords (which often require 2 CZs together with single-qubit gates once compiled), in addition to repeating these same circuits interleaved with a CZ at every depth in order to obtain the CZ gate fidelity estimate. In contrast, CAFE uses only 16 different circuits and 5 depths $0 \leq n \leq 8$ of CZs (with at most two additional CZ layers for state preparation and measurement). Moreover, the single-qubit gates used by IRB to create the random Clifford gates introduce unwanted error sources from decoherence and systematic errors, a problem that CAFE alleviates altogether by only repeating the cycle circuit of interest. Note that, unlike CAFE, RB additionally provides an estimate of the average gate

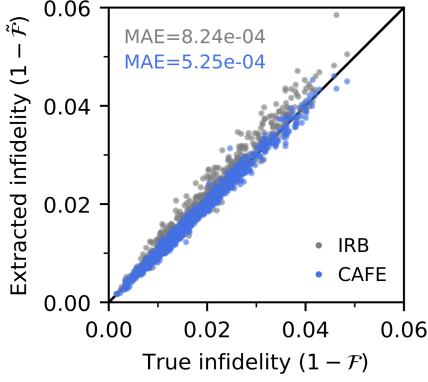


FIG. 2. Simulations comparing the CAFE and Interleaved Randomized Benchmarking (IRB) techniques for characterizing the average gate infidelity of noisy CZ gates with amplitude and phase damping noise in addition to coherent errors. For CAFE (blue), we use depths $n \in [0, 2, 4, 6, 8]$, whereas for IRB (gray) we use depths $n \in [5, 10, 15, 20, 25, 30, 35]$ for obtaining both the reference RB curve and the acquisition with interleaved CZ gates. Both approaches use 2000 shots per circuit. Labels presents the median absolute errors (MAE) of IRB and CAFE over $N = 1000$ different noisy CZ gates. $\text{MAE} = \text{median}(|\hat{\mathcal{F}}_1 - \mathcal{F}_1|, \dots, |\hat{\mathcal{F}}_N - \mathcal{F}_N|)$. We note that MAE scales with the true error of the gates being considered.

fidelity of all two-qubit Cliffords, which can be of independent interest. We also want to point out that using additional depths n and sampling shots can improve both approaches, and that considering a different range of noise parameters can change the resulting median absolute errors significantly. However, we have found that the conclusion remains that CAFE allows one to do a gate characterization at least as accurate as IRB, while requiring significantly less run time in experiment.

IV. BUDGETING COHERENT AND INCOHERENT ERRORS

Using CAFE, one can accurately estimate the fidelity of any unitary operation and report this single number as a performance metric, which is useful for validation and for estimating the performance of different quantum algorithms [1]. However, such a metric alone provides very little information as to how to improve the gate fidelity in practice. A central feature of CAFE is that one can extract actionable information about the origins of gate error by budgeting the coherent and incoherent contributions. To do so, one can fit the fidelity decay curve to a physical model, or modify the cycle circuit to echo out different parts of the gate errors, for example by leveraging Dynamical Decoupling (DD) pulses. We note that such budgeting is helpful for directing research focus, as interventions required to mitigate coherent and incoherent errors tend to be different.

A. Separating coherent and incoherent errors through fitting to a model

One way to obtain an error budget from CAFE is to fit the experimental data to a model. Although the CAFE experiment and resulting error budget are valid for any m -qubit unitary, we will again focus on the two-qubit CZ case to simplify the discussion in the main text. A more general derivation is presented in App. C. We assume that the cycle unitary is an excitation-preserving two-qubit gate close to a CZ

$$\tilde{U}(\Delta\theta, \Delta\gamma, \Delta\phi) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & e^{-i\Delta\gamma} \cos(\Delta\theta) & -ie^{-i\Delta\gamma} \sin(\Delta\theta) & 0 \\ 0 & -ie^{-i\Delta\gamma} \sin(\Delta\theta) & e^{-i\Delta\gamma} \cos(\Delta\theta) & 0 \\ 0 & 0 & 0 & -e^{-i(\Delta\phi+2\Delta\gamma)} \end{pmatrix}, \quad (2)$$

where $\tilde{U}(0, 0, 0) = \text{CZ}$, and the swap, single-qubit phase, and controlled-phased miscalibration angles are assumed to be small, such that $\Delta\theta, \Delta\gamma, \Delta\phi \ll 1$. To simplify further the expressions here, we also assume that the noisy quantum channel implementing the cycle circuit can be described by a two-qubit depolarizing channel

$$\mathcal{E}(\rho) = (1 - p_{\text{depol}}) \tilde{U} \rho \tilde{U}^\dagger + p_{\text{depol}} I_d / d, \quad (3)$$

which outputs a totally mixed state with probability p_{depol} and otherwise applies the cycle unitary $\tilde{U}$. With such a channel, the average gate fidelity for n repetitions of the cycle circuit is given by

$$\mathcal{F}_n = \frac{1}{4} - \epsilon_{\text{spam}} - \frac{1}{20} (1 - p_{\text{depol}})^n \cdot \left(1 - \left| 1 + 2e^{-in\Delta\gamma} \cos(n\Delta\theta) + e^{-in(2\Delta\gamma+\Delta\phi)} \right|^2 \right), \quad (4)$$

where we have explicitly included the SPAM errors ϵ_{spam} , which are assumed to vary slowly over the timescale of an experiment. We note that in this model we made a steady-state assumption about the gates in our system, but modeling transient behavior could be achieved with a more advanced fit. Using the expression in Eq. 4, we can model the experimental data to obtain the gate fidelity $\mathcal{F}$, in addition to the incoherent errors ϵ_{incoh} and coherent errors ϵ_{coh} of the quantum operation. To get these parameters in a way that is robust to SPAM errors, we use

$$1 - \mathcal{F} = 1 - \frac{\mathcal{F}_1}{(1 - \epsilon_{\text{spam}})}, \quad (5)$$

$$\epsilon_{\text{incoh}} = 1 - \frac{\mathcal{F}_1(p_{\text{depol}} = 0)}{(1 - \epsilon_{\text{spam}})}, \quad (6)$$

$$\epsilon_{\text{coh}} = 1 - \frac{\mathcal{F}_1(\Delta\theta = \Delta\gamma = \Delta\phi = 0)}{(1 - \epsilon_{\text{spam}})}. \quad (7)$$

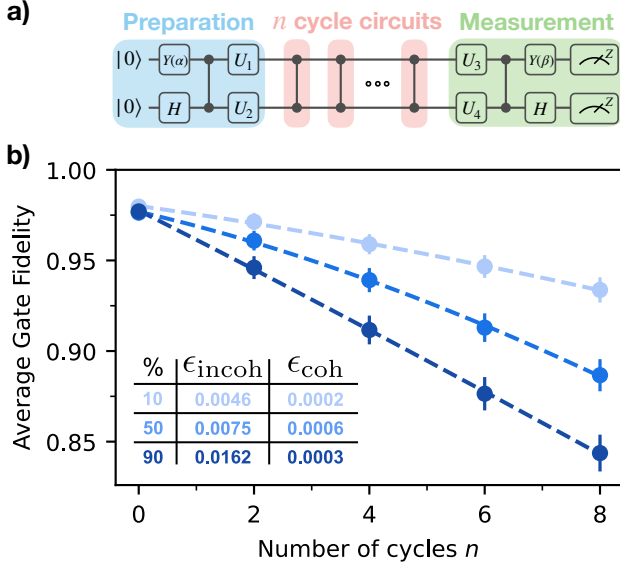


FIG. 3. Experimental characterization of CZ gates performed in parallel using CAFE. (a) Structure of the 16 circuits used to characterize the fidelity of a CZ gate for a given depth. The angles $\{\alpha, \beta\}$ and single-qubit unitaries U_j are found using the method described in App. B. Here the cycle circuit is a single CZ gate. (b) Data represents the cycle circuit fidelity compared to the ideal CZ unitary for different cycle repetitions, along with fits using the model presented in Eq. 4 and resulting gate budgets in the inset table. We plot the data for three specific gates – at the 10th, 50th and 90th percentile on a Sycamore device in terms of average gate infidelity – to illustrate the accuracy of the modeling over a wide range of data. Error bars represent the standard deviation of binomial distributions scaled by a factor of 5 to be visible.

In numerical simulations, which are presented in App. E, this analysis procedure was shown to be valid and robust to different unitary errors, as well as amplitude and phase damping channels. We use this method to budget errors in Figs. 2 to 5. A similar budgeting approach for single-qubit $X(\pi)$ gates, together with experimental results acquired on a Sycamore chip, are presented in App. D.

As shown in Fig. 3, the simple model of Eq. 4 allows us to accurately fit the CAFE curves spanning a wide range of CZ gate fidelities executed in parallel on a Sycamore chip. Additional data showing the consistency of the resulting error budget with XEB is presented in App. H.

A useful and intuitive picture to analyze the CAFE data is to consider the incoherent and coherent errors as linear and quadratic contributions to the gate infidelity, respectively. This can be seen directly by expanding Eq. 4 up to terms $\mathcal{O}((\Delta\theta)^4, (\Delta\gamma)^4, (\Delta\phi)^4, p_{\text{depol}}^2)$

$$\mathcal{F}_n \approx 1 - \epsilon_{\text{spam}} - \epsilon_{\text{lin}} n - \epsilon_{\text{quad}} n^2 \quad (8)$$

$$\epsilon_{\text{lin}} = \frac{3p_{\text{depol}}}{4} n \quad (9)$$

$$\epsilon_{\text{quad}} = \frac{8[(\Delta\theta)^2 + (\Delta\gamma)^2 + \Delta\gamma\Delta\phi] + 3(\Delta\phi)^2}{20} n^2. \quad (10)$$

This approximate quadratic form can be found for different cycle unitaries under similar noise channels, and could be used directly in the budgeting procedure given low gate infidelities and shallow cycle repetitions n , or in cases lacking an accurate analytical model for the quantum channel. Note that the fit model and resulting CAFE error budgeting can be straightforwardly extended to another platform or experiment. For example, the quantum channel describing the cycle circuit can be parameterized by a set of Kraus operators and computed numerically to fit the CAFE fidelities at different depths n . However, since CAFE performs only the minimal set of circuits needed to extract the average gate fidelity, the data will not contain enough information to fully characterize generic quantum channels, and additional experiments or other characterization tools should be considered if that were the goal.

We also note that the quantity ϵ_{incoh} estimates the average gate infidelity when no coherent control errors are present. This useful characterization metric is defined as $R(\mathcal{E})$ in Ref. [32], where the authors show that it bounds the *unitarity* of the channel $u(\mathcal{E})$:

$$\frac{d-1}{d}(1 - \sqrt{u(\mathcal{E})}) \leq R(\mathcal{E}). \quad (11)$$

As such, we can also relate our incoherent error estimate to unitarity. For instance, in the case of depolarizing noise, the unitarity is

$$u(\mathcal{E}_{\text{depol}}) = (1 - p_{\text{depol}})^2, \quad (12)$$

and the incoherent error obtained from CAFE is, as derived in App. C,

$$1 - \epsilon_{\text{incoh}} = \frac{1}{d} + \frac{d-1}{d}(1 - p_{\text{depol}}) \quad (13)$$

$$= \frac{1}{d} + \frac{d-1}{d}\sqrt{u(\mathcal{E}_{\text{depol}})}. \quad (14)$$

B. Isolating Coherent Channels using Dynamical Decoupling

Exploiting the versatility of CAFE, a second method for error budgeting is to modify the cycle circuit in order to isolate specific error channels. In particular, this is viable when the impact of the modifications is insignificant relative to the error channels being isolated. As a relevant example, we can leverage the fact that the cycle circuit is repeated n times by inserting dynamical decoupling gates in between repetitions. This approach of combining DD and CAFE, which we refer to as *DE-CAF*, is particularly useful to characterize the amount of certain coherent error present in the cycle circuit without necessitating full unitary tomography. We note that this is somewhat similar to the work in Ref. [5], with a single repetition of the cycle circuit, and the randomized unitaries replaced with specific DD pulses.

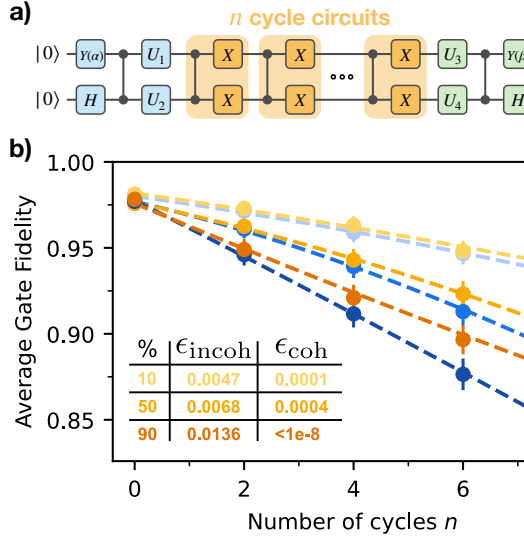


FIG. 4. CAFE combined with dynamical decoupling (DECAF). (a) Structure of the 16 DECAF circuits for characterizing a CZ gate at a given depth. Only the cycle circuit highlighted in orange differs from the CAFE circuits of Fig. 3a. (b) Experimental results on characterizing CZ gates in parallel when dynamical decoupling gates are part of the cycle circuit (orange), alongside the standard CAFE data presented in Fig. 3 (blue). Interleaving the CZ gates with X gates on both qubits echoes out low-frequency Z noise, which contributes to ϵ_{incoh} , and removes the sensitivity to single-qubit phase unitary errors, which contribute to ϵ_{coh} . Performing the simple DECAF experiment thus provides valuable information about the error mechanisms in CZ gates. Error bars are scaled by a factor 5 to be visible.

In the case of a CZ gate, adding an X gate to both qubits in the CAFE cycle circuit echos out the single-qubit phase errors, in addition to mitigating low-frequency noise, as demonstrated in App. F [33, 34]. As shown in Fig. 4, the DECAF data has higher fidelities and decreases more linearly than the CAFE data in practice. Looking at the resulting fits over all characterized CZ gates, we see a significant decrease in the median coherent error from 5.5×10^{-4} to 9.4×10^{-5} , which highlights single-qubit phase miscalibrations, alongside a decrease of the median incoherent error from 7.2×10^{-3} to 6.7×10^{-3} , which indicates the level of low-frequency noise in the device.

V. VALIDATING UNITARY CHARACTERIZATIONS

One key difference between CAFE and other methods for extracting error rates is the final unentangling step before measurement. By doing all of the inversion in a single step, similar to RB, but using an arbitrary characterized reference unitary for the inversion, we can validate different unitary characterizations while respecting the non-Clifford nature of most coherent error models. As

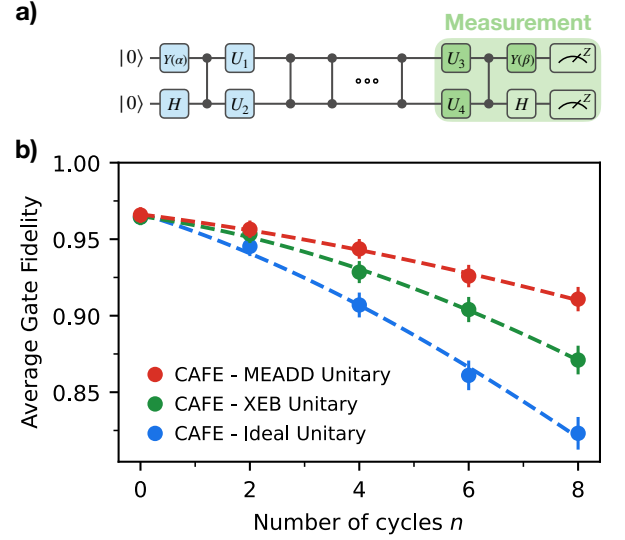


FIG. 5. Using CAFE to validate different unitary characterizations of a CZ gate with 50th percentile infidelity. (a) Only the three single-qubit gates $\{U_3, U_4, Y(\beta)\}$ highlighted in green need to be modified to compare the fidelity of the experimental CZ gate with an arbitrary unitary. (b) The different curves represent the fidelity between the experimental gate and an ideal CZ unitary (blue), a unitary extracted from an XEB experiment (green), and a unitary characterized with a newly introduced method called MEADD [35]. Since only the pre-measurement circuit changes in the three different CAFE experiments, the improvement in fidelity can entirely be attributed to considering a gate unitary that is closer to the experimental CZ implementation. Error bars are scaled by a factor 5 to be visible.

such, we can use CAFE to benchmark unitary characterizations, simply by changing the final measurement step and seeing which predictions most accurately map the final state back to $|0\rangle^{\otimes m}$, in conjunction with the methods presented in Sec. IV A. As illustrated in Fig. 5(a), for a two-qubit cycle circuit, this change corresponds to modifying three single-qubit gates in the last step of the CAFE protocol.

In Fig. 5(b), we compare the CZ unitary extracted by XEB to the one extracted using a unitary characterization method called Matrix-Element Amplification by Dynamical Decoupling (MEADD) [35], which is a characterization technique that we have developed based on Floquet characterization [36, 37]. MEADD allows us to isolate and precisely measure the different unitary parameters in any phased fSim gate (a general excitation number preserving two-qubit gate). We can see very clearly that the unitary predicted by MEADD is significantly better at predicting the coherent error than the unitary extracted by XEB. By increasing the amount of context around the gate, for example including some microwave operations or measurements on the surrounding qubits to include crosstalk or measurement-induced dephasing effects, we can see which characterizations break down in other contexts (see Sec. VI), and build trust that the

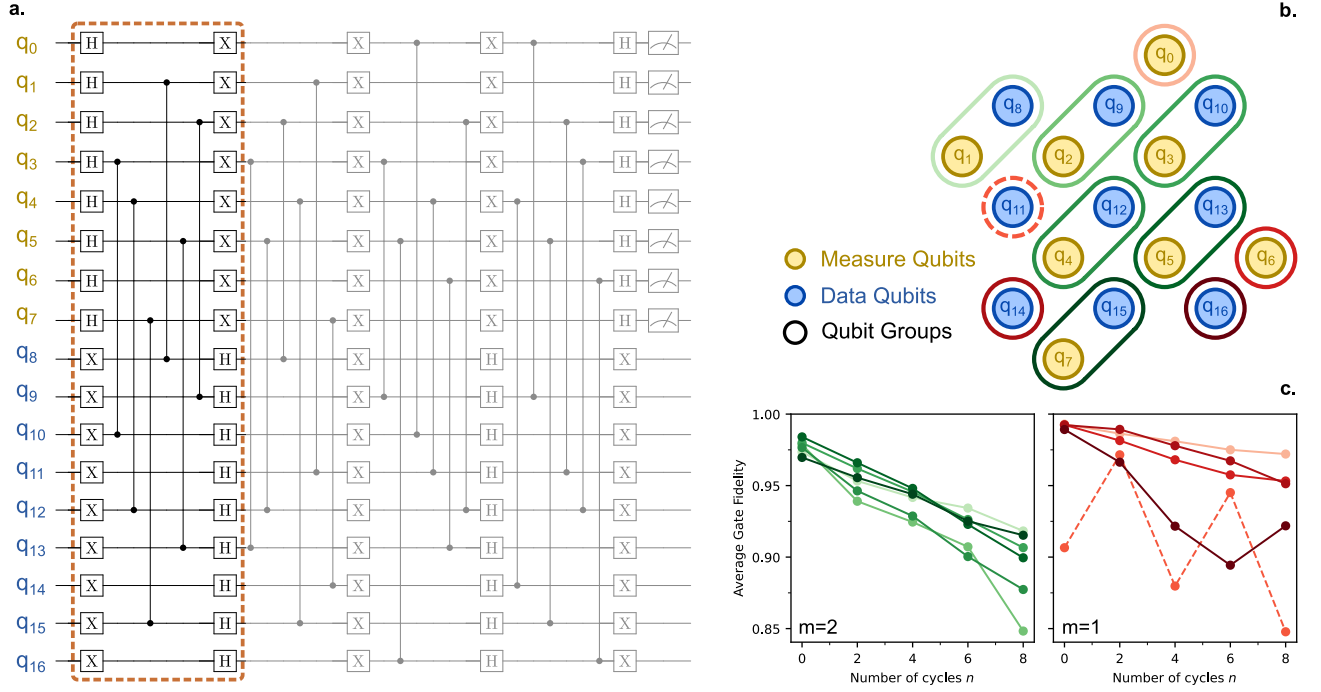


FIG. 6. **a.** The stabilizer extraction circuit for a distance-3 surface code, with the section to be characterized boxed in orange. **b.** A layout for the CAFE experiment along with the qubit groups characterized. **c.** CAFE data for repetitions of the cycle circuit from **a.** on the qubit groups in **b.**, with the $m = 2$ groups in green and the $m = 1$ groups in red. CAFE allows us to see that q_{11} (red, dotted) is experiencing a significant context-dependent error.

structures seen in our characterizations are the dominant effects impacting algorithm performance on the processor.

VI. CHARACTERIZING MULTI-LAYER CIRCUITS USING CAFE

Here, we provide an example of a CAFE experiment that includes both spatial and temporal context by characterizing a cycle circuit containing multiple layers. As illustrated in Fig. 6a, we consider a portion of the stabilizer extraction circuit for a distance-3 surface code experiment used in Ref. [26]. In Fig. 6b, we show how this section of the circuit can be considered on its own and broken down into qubit groups. We can then run a parallel CAFE experiment on these qubit groups which uses the highlighted section as the cycle circuit. The resulting experimental data is shown in Fig. 6c. We can see that one of the $m = 1$ qubit groups experiences significant error, despite its individual single-qubit gates X and H showing good fidelities when characterized in isolation. This highlights the importance of context-aware characterization, as the error mode being experienced was only visible when the gates were run in conjunction. This experiment also highlights the flexibility of CAFE in being able to characterize multi-operation circuits simultaneously on qubit groups of varying size.

VII. CONCLUSION

CAFE separates itself from other gate characterization methods due to its simplicity, efficiency, and flexibility, most notably as a complementary tool to other gate characterizations that provide more granular output. We have found its ability to split coherent and incoherent errors more reliable than other methods, and the ability to experimentally test unitary characterizations has allowed us to design and evaluate novel characterization methods more effectively.

Overall, our results demonstrate that CAFE yields accurate fidelity and error budget estimates using significantly less time in practice than complementary characterization protocols. This approach is thus particularly useful to quickly identify the worst performing individual quantum operations, typically single- and two-qubit gates, directly in the context of the larger quantum circuit in which they are found. Moreover, CAFE yields actionable information about the origins of the gate errors, namely the contributions of incoherent and coherent noise, which allow one to calibrate the quantum device to optimize performance on the experiment of interest. This usefulness of CAFE is demonstrated in Sec. VI for characterizing and re-calibrating a portion of the stabilizer extraction circuit of a distance-3 surface code experiment.

The flexibility of CAFE allows for many other as-of-yet unexplored variations, from creating fits which include

the impact of leakage, to changing the cycle circuit round-by-round to echo out components in different ways, to analyzing the different input bases separately to extract details about the error structure. Our hope is that other researchers will be able to create their own modifications of the CAFE framework to study the errors facing their own systems.

VIII. ACKNOWLEDGEMENTS

We are grateful to the Google Quantum AI team for building, operating, and maintaining software and hard-

ware infrastructure used in this work. The authors would like to thank Will Livingston, Sergio Boixo, Ben Chiaro, Vinicius S. Ferreira, Abraham Asfaw, and Paul V. Klimov for discussions and feedback on the draft, and Catherine Erickson and Andreas Bengtsson for software support.

-
- [1] J. Eisert, D. Hangleiter, N. Walk, I. Roth, D. Markham, R. Parekh, U. Chabaud, and E. Kashefi, “Quantum certification and benchmarking,” *Nature Reviews Physics* **2**, 382 (2020).
 - [2] E. Knill, D. Leibfried, R. Reichle, J. Britton, R. B. Blakestad, J. D. Jost, C. Langer, R. Ozeri, S. Seidelin, and D. J. Wineland, “Randomized benchmarking of quantum gates,” *Physical Review A* **77**, 012307 (2008).
 - [3] E. Magesan, J. M. Gambetta, and J. Emerson, “Scalable and robust randomized benchmarking of quantum processes,” *Physical review letters* **106**, 180504 (2011).
 - [4] J. J. Wallman and S. T. Flammia, “Randomized benchmarking with confidence,” *New Journal of Physics* **16**, 103032 (2014).
 - [5] S. Sheldon, L. S. Bishop, E. Magesan, S. Filipp, J. M. Chow, and J. M. Gambetta, “Characterizing errors on qubit operations via iterative randomized benchmarking,” *Physical Review A* **93**, 012301 (2016).
 - [6] T. Proctor, S. Seritan, K. Rudinger, E. Nielsen, R. Blume-Kohout, and K. Young, “Scalable randomized benchmarking of quantum computers using mirror circuits,” *Physical Review Letters* **129**, 150502 (2022).
 - [7] A. M. Pollreno, A. Carignan-Dugas, J. Hines, R. Blume-Kohout, K. Young, and T. Proctor, “A theory of direct randomized benchmarking,” (2023), arXiv:2302.13853 [quant-ph].
 - [8] S. Boixo, S. V. Isakov, V. N. Smelyanskiy, R. Babbush, N. Ding, Z. Jiang, M. J. Bremner, J. M. Martinis, and H. Neven, “Characterizing quantum supremacy in near-term devices,” *Nature Physics* **14**, 595 (2018).
 - [9] F. Arute, K. Arya, R. Babbush, D. Bacon, J. C. Bardin, R. Barends, R. Biswas, S. Boixo, F. G. Brandao, D. A. Buell, *et al.*, “Quantum supremacy using a programmable superconducting processor,” *Nature* **574**, 505 (2019).
 - [10] Y. Gu, W.-F. Zhuang, X. Chai, and D. E. Liu, “Benchmarking universal quantum gates via channel spectrum,” arXiv preprint arXiv:2301.02056 (2023).
 - [11] M. P. da Silva, O. Landon-Cardinal, and D. Poulin, “Practical Characterization of Quantum Devices without Tomography,” *Physical Review Letters* **107**, 210404 (2011), publisher: American Physical Society.
 - [12] S. T. Flammia and Y.-K. Liu, “Direct Fidelity Estimation from Few Pauli Measurements,” *Physical Review Letters* **106**, 230501 (2011), publisher: American Physical Society.
 - [13] A. Elben, S. T. Flammia, H.-Y. Huang, R. Kueng, J. Preskill, B. Vermersch, and P. Zoller, “The randomized measurement toolbox,” *Nature Reviews Physics* , 1 (2022), publisher: Nature Publishing Group.
 - [14] M. A. Nielsen and I. Chuang, “Quantum computation and quantum information,” (2002).
 - [15] C. H. Baldwin, A. Kalev, and I. H. Deutsch, “Quantum process tomography of unitary and near-unitary maps,” *Physical Review A* **90**, 012110 (2014).
 - [16] G. M. D’Ariano, M. G. Paris, and M. F. Sacchi, “Quantum tomography,” *Advances in Imaging and Electron Physics* **128**, 206 (2003).
 - [17] S. T. Merkel, J. M. Gambetta, J. A. Smolin, S. Poletto, A. D. Córcoles, B. R. Johnson, C. A. Ryan, and M. Steffen, “Self-consistent quantum process tomography,” *Physical Review A* **87**, 062119 (2013).
 - [18] D. Greenbaum, “Introduction to quantum gate set tomography,” arXiv preprint arXiv:1509.02921 (2015).
 - [19] E. Nielsen, J. K. Gamble, K. Rudinger, T. Scholten, K. Young, and R. Blume-Kohout, “Gate set tomography,” *Quantum* **5**, 557 (2021).
 - [20] J. Kelly, P. O’Malley, M. Neeley, H. Neven, and J. M. Martinis, “Physical qubit calibration on a directed acyclic graph,” arXiv:1803.03226 [quant-ph] (2018), arXiv: 1803.03226.
 - [21] K. Rudinger, T. Proctor, D. Langharst, M. Sarovar, K. Young, and R. Blume-Kohout, “Probing Context-Dependent Errors in Quantum Processors,” *Physical Review X* **9**, 021045 (2019), publisher: American Physical Society.
 - [22] Y. Wu, W.-S. Bao, S. Cao, F. Chen, M.-C. Chen, X. Chen, T.-H. Chung, H. Deng, Y. Du, D. Fan, *et al.*, “Strong quantum computational advantage using a superconducting quantum processor,” *Physical review letters* **127**, 180501 (2021).
 - [23] P. Murali, D. C. McKay, M. Martonosi, and A. Javadi-Abhari, “Software Mitigation of Crosstalk on Noisy Intermediate-Scale Quantum Computers,” in *Proceedings of the Twenty-Fifth International Conference on Architectural Support for Programming Languages and Operating Systems*, ASPLOS ’20 (Association for Computing Machinery, New York, NY, USA, 2020) pp. 1001–1016.
 - [24] M. Sarovar, T. Proctor, K. Rudinger, K. Young, E. Nielsen, and R. Blume-Kohout, “Detecting crosstalk

- errors in quantum information processors,” *Quantum* **4**, 321 (2020).
- [25] S. Krinner, N. Lacroix, A. Remm, A. Di Paolo, E. Genois, C. Leroux, C. Hellings, S. Lazar, F. Swiadek, J. Herrmann, *et al.*, “Realizing repeated quantum error correction in a distance-three surface code,” *Nature* **605**, 669 (2022).
 - [26] G. Q. AI, “Suppressing quantum errors by scaling a surface code logical qubit,” (2022).
 - [27] C. Ryan-Anderson, J. G. Bohnet, K. Lee, D. Gresh, A. Hankin, J. P. Gaebler, D. Francois, A. Chernoguzov, D. Lucchetti, N. C. Brown, T. M. Gatterman, S. K. Halit, K. Gilmore, J. A. Gerber, B. Neyenhuis, D. Hayes, and R. P. Stutz, “Realization of real-time fault-tolerant quantum error correction,” *Phys. Rev. X* **11**, 041058 (2021).
 - [28] C. Dankert, R. Cleve, J. Emerson, and E. Livine, “Exact and approximate unitary 2-designs and their application to fidelity estimation,” *Physical Review A* **80**, 012304 (2009).
 - [29] A. W. Harrow and R. A. Low, “Random quantum circuits are approximate 2-designs,” *Communications in Mathematical Physics* **291**, 257 (2009).
 - [30] H. Zhu, Y. S. Teo, and B.-G. Englert, “Two-qubit symmetric informationally complete positive-operator-valued measures,” *Physical Review A* **82**, 042308 (2010).
 - [31] D. Lu, H. Li, D.-A. Trottier, J. Li, A. Brodutch, A. P. Krismanich, A. Ghavami, G. I. Dmitrienko, G. Long, J. Baugh, and R. Laflamme, “Experimental Estimation of Average Fidelity of a Clifford Gate on a 7-Qubit Quantum Processor,” *Physical Review Letters* **114**, 140505 (2015), publisher: American Physical Society.
 - [32] J. Wallman, C. Granade, R. Harper, and S. T. Flammia, “Estimating the coherence of noise,” *New Journal of Physics* **17**, 113020 (2015).
 - [33] L. Viola and S. Lloyd, “Dynamical suppression of decoherence in two-state quantum systems,” *Physical Review A* **58**, 2733 (1998).
 - [34] L. Viola and E. Knill, “Robust dynamical decoupling of quantum systems with bounded controls,” *Physical Review Letters* **90**, 037901 (2003).
 - [35] E. Genois, J. A. Gross, Z.-P. Cian, D. M. Debroy, and Z. Jiang, “Matrix element amplification using dynamical decoupling,” in preparation.
 - [36] F. Arute, K. Arya, R. Babbush, D. Bacon, J. C. Bardin, R. Barends, A. Bengtsson, S. Boixo, M. Broughton, B. B. Buckley, D. A. Buell, B. Burkett, N. Bushnell, Y. Chen, Z. Chen, Y.-A. Chen, B. Chiaro, R. Collins, S. J. Cotton, W. Courtney, S. Demura, A. Derk, A. Dunsworth, D. Eppens, T. Ekl, C. Erickson, E. Farhi, A. Fowler, B. Foxen, C. Gidney, M. Giustina, R. Graff, J. A. Gross, S. Habegger, M. P. Harrigan, A. Ho, S. Hong, T. Huang, W. J. Huggins, L. B. Ioffe, S. V. Isakov, E. Jeffrey, Z. Jiang, C. Jones, D. Kafri, K. Kechedzhi, J. Kelly, S. Kim, P. V. Klimov, A. N. Korotkov, F. Kostritsa, D. Landhuis, P. Laptev, M. Lindmark, E. Lucero, M. Marthaler, O. Martin, J. M. Martinis, A. Marusczyk, S. McArdle, J. R. McClean, T. McCourt, M. McEwen, A. Megrant, C. Mejuto-Zaera, X. Mi, M. Mohseni, W. Mruczkiewicz, J. Mutus, O. Naaman, M. Neeley, C. Neill, H. Neven, M. Newman, M. Y. Niu, T. E. O’Brien, E. Ostby, B. Pató, A. Petukhov, H. Putterman, C. Quintana, J.-M. Reiner, P. Roushan, N. C. Rubin, D. Sank, K. J. Satzinger, V. Smelyanskiy, D. Strain, K. J. Sung, P. Schmitteckert, M. Szalay, N. M. Tubman, A. Vainsencher, T. White, N. Vogt, Z. J. Yao, P. Yeh, A. Zalcman, and S. Zanker, “Observation of separated dynamics of charge and spin in the Fermi-Hubbard model,” *arXiv:2010.07965* (2020).
 - [37] C. Neill, T. McCourt, X. Mi, Z. Jiang, M. Y. Niu, W. Mruczkiewicz, I. Aleiner, F. Arute, K. Arya, J. Atalaya, R. Babbush, J. C. Bardin, R. Barends, A. Bengtsson, A. Bourassa, M. Broughton, B. B. Buckley, D. A. Buell, B. Burkett, N. Bushnell, J. Campero, Z. Chen, B. Chiaro, R. Collins, W. Courtney, S. Demura, A. R. Derk, A. Dunsworth, D. Eppens, C. Erickson, E. Farhi, A. G. Fowler, B. Foxen, C. Gidney, M. Giustina, J. A. Gross, M. P. Harrigan, S. D. Harrington, J. Hilton, A. Ho, S. Hong, T. Huang, W. J. Huggins, S. V. Isakov, M. Jacob-Mitos, E. Jeffrey, C. Jones, D. Kafri, K. Kechedzhi, J. Kelly, S. Kim, P. V. Klimov, A. N. Korotkov, F. Kostritsa, D. Landhuis, P. Laptev, E. Lucero, O. Martin, J. R. McClean, M. McEwen, A. Megrant, K. C. Miao, M. Mohseni, J. Mutus, O. Naaman, M. Neeley, M. Newman, T. E. O’Brien, A. Opremcak, E. Ostby, B. Pató, A. Petukhov, C. Quintana, N. Redd, N. C. Rubin, D. Sank, K. J. Satzinger, V. Shvarts, D. Strain, M. Szalay, M. D. Trevithick, B. Villalonga, T. C. White, Z. Yao, P. Yeh, A. Zalcman, H. Neven, S. Boixo, L. B. Ioffe, P. Roushan, Y. Chen, and V. Smelyanskiy, “Accurately computing the electronic properties of a quantum ring,” *Nature* **594**, 508 (2021).
 - [38] F. Yan, P. Krantz, Y. Sung, M. Kjaergaard, D. L. Campbell, T. P. Orlando, S. Gustavsson, and W. D. Oliver, “Tunable coupling scheme for implementing high-fidelity two-qubit gates,” *Phys. Rev. Appl.* **10**, 054062 (2018).
 - [39] “Exponential suppression of bit or phase errors with cyclic error correction,” *Nature* **595**, 383 (2021).
 - [40] C. Developers, “Cirq,” (2022), See full list of authors on Github: <https://github.com/quantumlib/Cirq/graphs/contributors>.

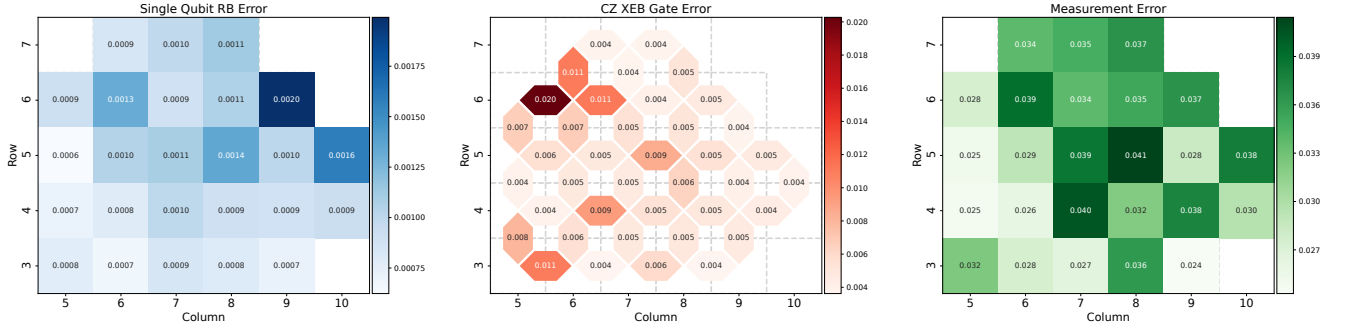


FIG. 7. Heatmaps for per-qubit and per-pair median error rates for single-qubit Randomized Benchmarking, CZ XEB, and average readout identification error.

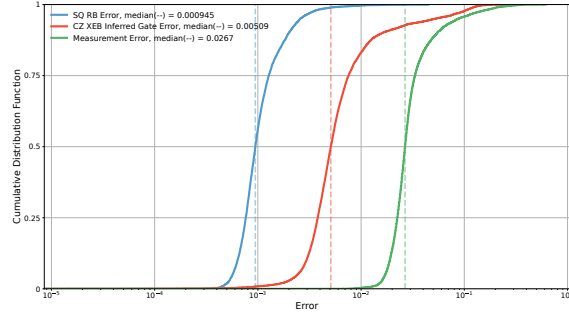


FIG. 8. Cumulative density functions for single-qubit RB, CZ XEB, and readout errors.

Appendix A: Device Specifics

All experimental data presented here was collected on a subset of a Sycamore device with 72 transmon qubits and 121 tunable couplers, where each qubit is coupled to up to four neighbors. In Fig. 7 we present median error rates for microwave single qubit gates, CZ gates, and measurements over the timeframe data was acquired for this paper. All errors reported are measured simultaneously. For more information on the architecture see Refs. [9] and [38], and for more information on gates implementation see Ref. [39].

Appendix B: Preparing and measuring two-qubit states

We describe a method to prepare or measure an arbitrary two-qubit state with a single maximally entangling operation. We start by considering a general two-qubit target state

$$|\psi\rangle = A|00\rangle + B|01\rangle + C|10\rangle + D|11\rangle, \quad (B1)$$

$$|A|^2 + |B|^2 + |C|^2 + |D|^2 = 1.$$

We can then write a matrix with the amplitudes of this state

$$M_\psi = \begin{bmatrix} A & B \\ C & D \end{bmatrix}, \quad (B2)$$

and perform the singular value decomposition (SVD) $M_\psi = U_\psi S_\psi V_\psi^\dagger$. Considering the initial singular value, S_ψ^{00} , we can quantify the level of entanglement in the targeted state. The first step to generate this state is to prepare a state with a matching entanglement signature.

In the case where we intend to use a CZ gate, this can be done by putting one qubit in $|+\rangle$, and applying a Y rotation to the other qubit with an angle of $\alpha = 2\arccos(S_\psi^{00})$. This state, labeled as $|CZ\rangle$ in Fig. 9, is equivalent to the final

state up to single qubit rotations. To find these rotations, we take the SVD of the 2×2 matrix corresponding to this intermediate state, $U_{CZ} S_{CZ} V_{CZ}^\dagger$. The single qubit unitary required for the first qubit is given by $U_1 = U_\psi U_{CZ}^{-1}$, and the unitary for the second qubit is $U_2 = V_\psi V_{CZ}^{-1}$. The resulting circuit is shown in Fig. 9. Constructions for CZ and $\sqrt{\text{iSWAP}}$ are implemented as the methods `prepare_two_qubit_state_with_cz` and `prepare_two_qubit_state_with_iswap` in the open-source software `Cirq` [40].

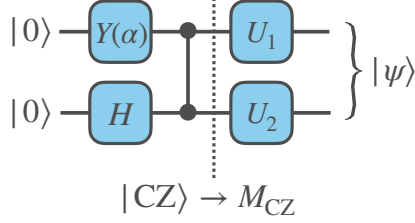


FIG. 9. A simple circuit that maps $|00\rangle$ to an arbitrary two-qubit state $|\psi\rangle$ using only one CZ operation. The overlap of any two-qubit state with the state $|\psi\rangle$ can be obtained by executing the inverse of this circuit and reporting the probability of measuring $|00\rangle$ afterwards.

Appendix C: Derivation of Eq. 4

In this section, we derive Eq. 4 as an example of how one could do the same for other unitaries and error models of interest. Given a general quantum channel described by a set of Kraus operators K_α ,

$$\mathcal{E}(\rho) = \sum_{\alpha} K_{\alpha} \rho K_{\alpha}^{\dagger}, \quad (\text{C1})$$

the average fidelity of the quantum operation $\mathcal{E}$ with respect to a unitary operation U is

$$\mathcal{F}(\mathcal{E}, U) = \int \langle \psi(\mathbf{x}) | U^{\dagger} \mathcal{E}(|\psi(\mathbf{x})\rangle \langle \psi(\mathbf{x})|) U | \psi(\mathbf{x}) \rangle d\mu(\mathbf{x}) \quad (\text{C2})$$

$$= \sum_{\alpha} \int |\langle \psi(\mathbf{x}) | U^{\dagger} K_{\alpha} |\psi(\mathbf{x})\rangle|^2 d\mu(\mathbf{x}), \quad (\text{C3})$$

where $\mathbf{x}$ is a parametrization of pure quantum states and $\mu(\mathbf{x})$ is the Haar measure. To evaluate $\mathcal{F}$, we introduce the projector on the symmetric subspace of the system and its replica, which is expressed as

$$P_S = \frac{d(d+1)}{2} \int |\psi(\mathbf{x})\rangle \langle \psi(\mathbf{x})| \otimes |\psi(\mathbf{x})\rangle \langle \psi(\mathbf{x})| d\mu(\mathbf{x}), \quad (\text{C4})$$

where $d = 2^m$ is the dimension of the m -qubit Hilbert space. Using this projector, the fidelity takes the form

$$\mathcal{F}(\mathcal{E}, U) = \frac{2}{d(d+1)} \text{tr} \left(P_S \sum_{\alpha} K_{\alpha}^{\dagger} U \otimes U^{\dagger} K_{\alpha} \right) \quad (\text{C5})$$

$$= \frac{1}{d(d+1)} \sum_{\alpha} \text{tr} K_{\alpha}^{\dagger} K_{\alpha} + |\text{tr}(U^{\dagger} K_{\alpha})|^2 \quad (\text{C6})$$

$$= \frac{1}{d+1} + \frac{1}{d(d+1)} \sum_{\alpha} |\text{tr}(U^{\dagger} K_{\alpha})|^2. \quad (\text{C7})$$

In order to compute the fidelity expression in Eq. C7, one computes the eigenvalues of $U^{\dagger} K_{\alpha}$ and sum them up directly to get $\text{tr}(U^{\dagger} K_{\alpha})$. For example, we consider the case where the operation of interest can be described by a quantum channel that either applies the unitary $\tilde{U}$, or totally depolarizes the qubits with probability p_{depol}

$$\mathcal{E}(\rho) = (1 - p_{\text{depol}}) \tilde{U} \rho \tilde{U}^{\dagger} + p_{\text{depol}} I_d/d, \quad (\text{C8})$$

where I_d is the d -dimensional identity operator. The fidelity of this channel after n consecutive applications is given by

$$\mathcal{F}_n = (1 - p_{\text{depol}})^n \frac{d + |\text{tr}[(U^\dagger)^n \tilde{U}^n]|^2}{d(d+1)} + [1 - (1 - p_{\text{depol}})^n] \frac{1}{d}. \quad (\text{C9})$$

After subtracting ϵ_{spam} from the RHS to account for experimental SPAM errors, the expression in Eq. C9 can be used to fit the CAFE data of any m -qubit unitary $\tilde{U}$, subject to symmetric depolarizing noise, with respect to a reference unitary U . Note that in the case where the quantum operation under characterization, $\tilde{U}$, corresponds exactly with the reference unitary U , the CAFE data should obey the following single-exponential decay

$$\mathcal{F}_n = \frac{1}{d} + \frac{d-1}{d}(1 - p_{\text{depol}})^n. \quad (\text{C10})$$

Otherwise, when $\tilde{U} \neq U$, coherent errors are present and the first term of C9 will introduce errors that scale quadratically with n to first order.

In the case we are interested in, we want to use CAFE to characterize a two-qubit gate, where $d = 4$, $U = \text{CZ}$ and $\tilde{U}$ is a number-preserving fSim gate close to a CZ gate:

$$\tilde{U}(\Delta\theta, \Delta\gamma, \Delta\phi) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & e^{-i\Delta\gamma} \cos(\Delta\theta) & -ie^{-i\Delta\gamma} \sin(\Delta\theta) & 0 \\ 0 & -ie^{-i\Delta\gamma} \sin(\Delta\theta) & e^{-i\Delta\gamma} \cos(\Delta\theta) & 0 \\ 0 & 0 & 0 & -e^{-i(2\Delta\gamma+\Delta\phi)} \end{pmatrix}, \quad (\text{C11})$$

with some residual SWAP-like error captured by the angle $\Delta\theta$, single-qubit phase error captured by $\Delta\gamma$, and CPhase-like error captured by $\Delta\phi$, with $\Delta\theta, \Delta\gamma, \Delta\phi \ll 1$. Since CZ is diagonal, we can show that the eigenvalues of $(U^\dagger)^n \tilde{U}^n$ are

$$\boldsymbol{\lambda} = (1 \ e^{-in(\Delta\gamma-\Delta\theta)} \ e^{-in(\Delta\gamma+\Delta\theta)} \ e^{-in(2\Delta\gamma+\Delta\phi)}), \quad (\text{C12})$$

for $n \in \mathbb{N}$. As such, $\text{tr}[(U^\dagger)^n \tilde{U}^n] = \text{sum}(\boldsymbol{\lambda})$ and substituting into Eq. C9 gives

$$\mathcal{F}_n = \frac{1}{4} - (1 - p_{\text{depol}})^n \left(\frac{1 - |1 + 2e^{-in\Delta\gamma} \cos(n\Delta\theta) + e^{-in(2\Delta\gamma+\Delta\phi)}|^2}{20} \right) \quad (\text{C13})$$

$$= 1 - \frac{3p_{\text{depol}}}{4}n - \frac{8(\Delta\theta)^2 + 8(\Delta\gamma)^2 + 8(\Delta\gamma\Delta\theta) + 3(\Delta\phi)^2}{20}n^2 \quad (\text{C14})$$

$$+ \mathcal{O}((\Delta\theta)^4, (\Delta\gamma)^4, (\Delta\phi)^4, p_{\text{depol}}^2), \quad (\text{C15})$$

which is the expression in Eq. 4 of the main text, after subtracting the SPAM errors ϵ_{spam} from the RHS to account for experimental imperfections.

Appendix D: Single-qubit CAFE

In order to give an additional example on how to deploy the CAFE characterization framework, we present an experiment characterizing 68 single-qubit X gates in parallel on a Sycamore processor. The results are presented in Fig. 10, showing the fidelity of the experimental gate with both an ideal unitary (in blue) and a characterized one (in red), in addition to fits using an analytical model derived below, for three example gates. In the single qubit case, the minimal 2-design contains only 4 states which can all be prepared with a single gate, which significantly speeds up the characterization. We have used cycle repetitions n up to 32 here, showing how one can modify the experiment when studying especially high-fidelity operations to maintain good fit robustness.

Assuming that the incoherent noise of the experimental gate can be described by a totally depolarizing channel, our model to fit the single-qubit CAFE data can again be derived from Eq. (C9). For this case where $U = R_x(\pi)$ and $d = 2$, we can parametrize a single-qubit unitary

$$\tilde{U}(\Delta\mu) = \begin{pmatrix} -\sin(\Delta\mu/2) & -i\cos(\Delta\mu/2) \\ -i\cos(\Delta\mu/2) & -\sin(\Delta\mu/2) \end{pmatrix}, \quad (\text{D1})$$

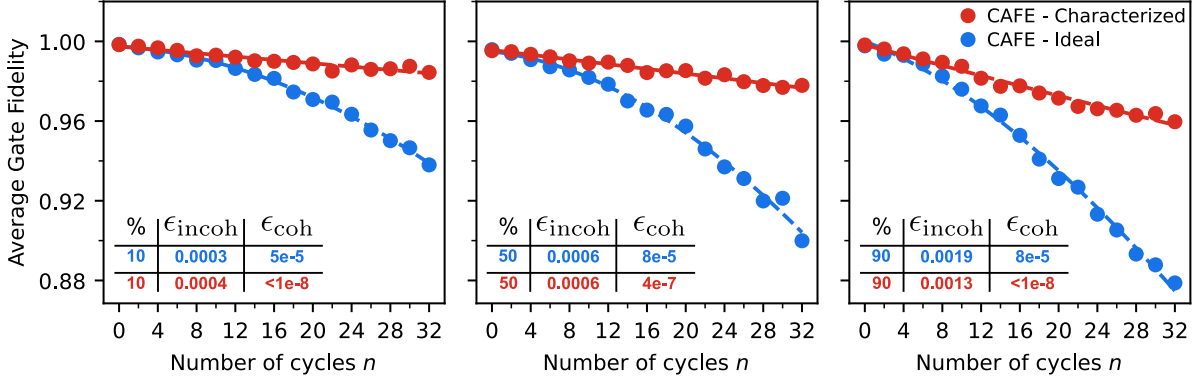


FIG. 10. Single-qubit parallel CAFE results for three X gates close to the 10th (left), 50th (center), and 90th (right) percentiles of infidelity, from a total of 68 parallel X gates on a Sycamore device. The data in blue computes the fidelity with an ideal $U = X$ unitary, whereas the CAFE data in red uses a characterized unitary. Note that in this data, the characterized single-qubit unitaries were not used to optimize the gate parameters, thus the imperfect calibration.

where $\tilde{U}(0) = R_x(\pi)$. One can easily show that the eigenvalues of $(U^\dagger)^n \tilde{U}^n$ are

$$\lambda_n = \begin{cases} (i^n e^{in\Delta\mu/2} & i^n e^{-in\Delta\mu/2}), \text{ for } n \text{ even} \\ (e^{in\Delta\mu/2} & e^{-in\Delta\mu/2}), \text{ for } n \text{ odd} \end{cases} \quad (\text{D2})$$

where $n \in \mathbb{N}$. As such,

$$\left| \text{tr}[(U^\dagger)^n \tilde{U}^n] \right|^2 = |\text{sum}(\lambda_n)|^2 = 4 \cos^2(n\Delta\mu/2) \quad (\text{D3})$$

and Eq. C9 becomes

$$\mathcal{F}_n = \frac{1}{2} - \epsilon_{\text{spam}} + (1 - p_{\text{depol}})^n \left(\frac{2}{3} \cos^2(n\Delta\mu/2) - \frac{1}{6} \right). \quad (\text{D4})$$

where, as before, we have explicitly included the SPAM errors ϵ_{spam} to account for experimental imperfections. The error budget in terms of the fidelity $\mathcal{F}$, incoherent error contribution ϵ_{incoh} , and coherent error contribution ϵ_{coh} is obtained the same way as in Eqs. (5) to (7) of the main text, i.e. by fitting the data with this expression, evaluating it at $n = 1$ with different noise parameters set to zero, and normalizing by the SPAM errors. Figure 10 shows that such a model allows us to accurately reproduce the CAFE data obtained in the experiment.

Appendix E: Numerical CAFE simulations

In this section, we present numerical simulations where we characterize a CZ gate using the CAFE approach presented in the main text. We simulate the same $2^4 = 16$ circuits executed on the Sycamore device with noisy two-qubit unitaries, together with either depolarizing noise in App. E1 or amplitude and phase damping noise in App. E2. We also used the same low cycle repetitions $n \in [0, 2, 4, 6, 8]$ which make the CAFE experiment have especially low execution time relative to other two-qubit benchmarking techniques. In Sec. III we showed that using CAFE to characterize a CZ gate can allow for a more accurate fidelity estimation than the widely used Randomized Benchmarking (RB) protocol while using significantly less experimental resources. Note that we also simulate the sampling of 2000 shots used experimentally, which places an upper bound on the standard deviation of all the CAFE data points presented in Figs. 3, 4, and 5 of

$$\sigma_{\mathcal{F}} \leq \frac{1}{8\sqrt{N_{\text{shots}}}} \approx 0.0028, \quad (\text{E1})$$

which is smaller than all the markers used in the plots of the main text.

Importantly, these simulations allow us to confirm the validity of the CAFE experiment and following fitting procedure to obtain accurate estimates of the following: the average gate fidelity of the quantum operation with regards to any reference unitary, the incoherent error contribution to the infidelity and the coherent error contribution.

The simulations were performed using the open-source software `Cirq` [40].

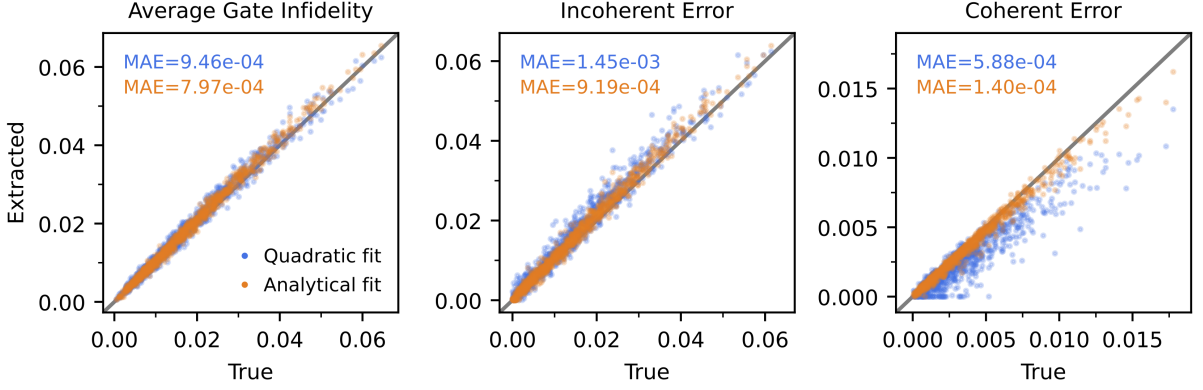


FIG. 11. Simulating the CAFE experiment for characterizing 1000 different CZ gates with realistically small coherent errors and depolarizing incoherent noise. Two error budgeting techniques obtained by fitting this CAFE data are presented. In blue, using a quadratic fit for the fidelity at depths $n \in [0, 2, 4]$, and in orange using the analytical expression of Eq. (4) for depths $n \in [0, 2, 4, 6, 8]$. Inset shows the median absolute errors (MAE) of the two error budgeting techniques from the true values inserted into the simulation. We note that MAE scales with the true error of the gates being considered. The true incoherent (coherent) errors are obtained from the channel's average gate infidelity when only depolarizing errors (unitary errors) are present in the simulation.

1. Depolarizing noise

In a first case, we consider that the noisy CZ gate we are trying to characterize is described by the quantum channel

$$\mathcal{E}(\rho) = (1 - p_{\text{depol}}) \tilde{V} \rho \tilde{V}^\dagger + p_{\text{depol}} I_d/d, \quad (\text{E2})$$

which outputs a totally depolarized state with probability p_{depol} , and otherwise applies the excitation-preserving unitary

$$\tilde{V}(\theta, \zeta, \chi, \gamma, \phi) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & e^{-i(\gamma+\zeta)} \cos \theta & -i e^{-i(\gamma-\chi)} \sin \theta & 0 \\ 0 & -i e^{-i(\gamma+\chi)} \sin \theta & e^{-i(\gamma-\zeta)} \cos \theta & 0 \\ 0 & 0 & 0 & -e^{-i(2\gamma+\phi)} \end{pmatrix}, \quad (\text{E3})$$

where we allow for miscalibrations in all of the five angles that parametrize this gate, where $\tilde{V}(0, 0, 0, 0, 0) = \text{CZ}$. For the simulation, we then sample the depolarizing probability uniformly in the range from 0 to 0.05, and sample the miscalibrated angles from a normal distribution with zero mean and standard deviation of 0.05 rad such that

$$p_{\text{depol}} \propto \mathcal{U}(0, 0.05) \quad (\text{E4})$$

$$\{\theta, \zeta, \chi, \gamma, \phi\} \propto \mathcal{N}(0, 0.05^2). \quad (\text{E5})$$

Using this model, we can simulate the noisy CAFE experiment for different cycle repetitions n , and then fit these data points to obtain an error budgeting of the noisy CZ gate, which we can compare directly with the true parameter values that were drawn for the simulation.

In Fig. 11, we present the average gate infidelity $(1 - \mathcal{F})$, incoherent error contribution ϵ_{incoh} , and coherent error contribution ϵ_{coh} to this infidelity for 1000 independently sampled unitaries and depolarizing probabilities. We present two valid gate error budgeting approaches, one which fits the CAFE data with the analytical expression in Eq. 4 (orange points), and one which uses the simple quadratic form of Eq. 8 (blue points). These scatter demonstrate the validity of CAFE to characterize the fidelity of a quantum operation and budget its coherent and incoherent contributions with an imprecision smaller or equal to 0.001 on median for realistic gate fidelities. Note that the reported median absolute errors depend on the specific range of gate fidelities considered. For example, the MAE values decrease significantly when considering only gate infidelities $< 1\%$ in the ensemble.

The quadratic fit results shown in blue in Fig. 11 are an indication of how the CAFE framework can be deployed in order to estimate the fidelity of an operation in context. However, this simplest approach gives slightly biased results for ϵ_{incoh} and ϵ_{coh} , which is understood from the breakdown of the $np_{\text{depol}} \ll 1$ approximation and/or the $n\alpha \ll 1$

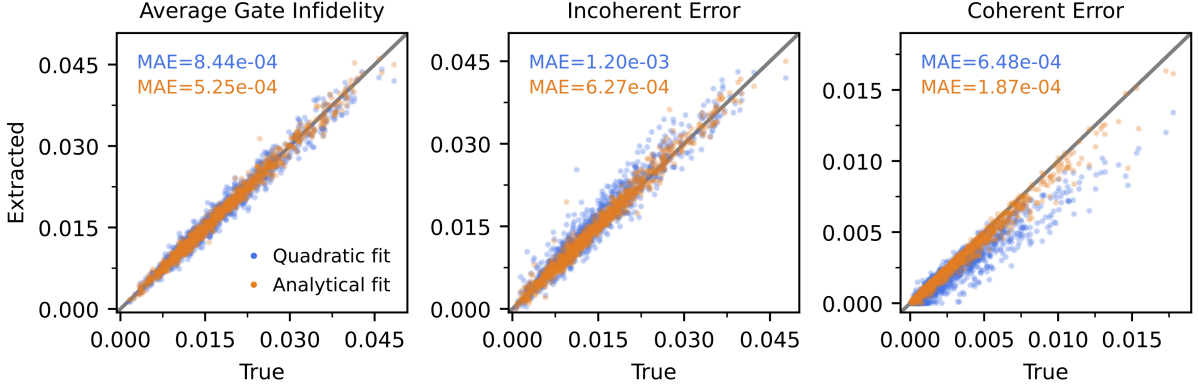


FIG. 12. Similar simulations as presented in Fig. 11, now implementing realistic amplitude and phase damping incoherent noise instead of a depolarizing channel. The analytical fit, which still uses Eq. (4) that effectively lumps the incoherent noise into a p_{depol} probability, works remarkably well even under this different noise channel.

approximation, where α here stands for any of the five miscalibrated CZ angles. Consequently, using a quadratic fit should be used carefully, for operations with high fidelity ($> 99\%$) and using shallow cycle repetitions n . In fact, for these simulations, we have found that using larger depths than $n = 4$ for the quadratic fit decreased the accuracy of the resulting error budgets. However, if the fidelity of the operation of interest is very high, for example with single-qubit gates, using a quadratic fit becomes an attractive strategy when lacking a proper analytical model of the error origins.

On the other hand, the analytical fit performs well for larger depths n since it does not rely on any approximation. Here, we have only used the depths $n \in [0, 2, 4, 6, 8]$ to be consistent with the experiment results of the main text. These results are an important validation for CAFE, but not necessarily surprising since we are using the same error model to simulate the experiment and to fit the resulting data. In the next sub-section, we show that this model also holds very well for different noise models such as amplitude and phase damping, provided they cause realistic gate infidelities of a few percent or less.

2. Amplitude and phase damping noise

To verify the robustness of the CAFE framework to different incoherent noise processes, we have also simulated a two-qubit CAFE experiment where the quantum channel is described by the same unitary $\tilde{V}$ with randomly sampled angles as before, but now followed by a single-qubit amplitude and phase damping channel on both qubits. The total decay and phase-flip probabilities for both qubits were sampled from a normal distribution with a standard deviation of 0.03, before taking the absolute value

$$\{p_{\text{decay,total}}, p_{\text{phaseflip,total}}\} \propto |\mathcal{N}(0, 0.03^2)|, \quad (\text{E6})$$

and the individual decay and phase-flip probabilities of the two qubits were splitting these total probabilities with a ratio drawn from a normal distribution with mean of 0.5 and standard deviation of 0.1, such that the qubits have similar but distinct coherence properties.

In Fig. 12, we show that the CAFE framework we developed again gives very accurate CZ characterizations and error budgets. In particular, using the analytical expression of Eq. (4) to fit this CAFE data works remarkably well given that the noise present in the gate cannot be fully described by the assumed depolarizing channel. Since the incoherent noise is small, but in the realistic range of creating about 0.1 % to 4 % gate infidelity, the fit is able to approximate well the incoherent error contribution into the single exponential of the model. This gives us confidence that we can leverage this model in realistic gate characterization scenarios. Moreover, the simulations show that the accuracy of using this model increases with increasing gate fidelities, which is a great indication for the continued usefulness of this method as gates and qubit coherences keep improving.

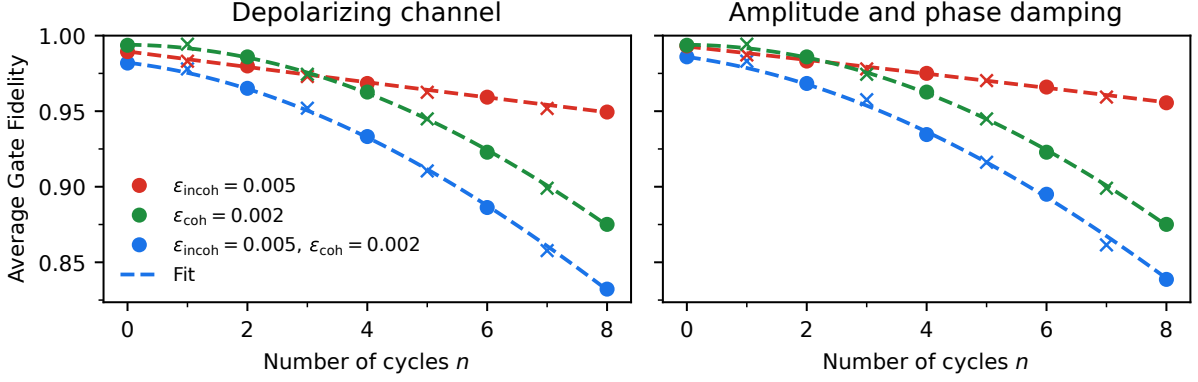


FIG. 13. Example of fitting the CAFE curve obtained in simulations with the depolarizing noise channel used in App. E 1 (left) and the amplitude and phase damping noise channel used in App. E 2 (right). In blue, we show simulations including both incoherent errors of strength $\epsilon_{\text{incoh}} = 0.005$ and coherent errors of $\epsilon_{\text{coh}} = 0.002$, and the dotted line shows the fit of this data using Eq. (4). As in the main text, only the even depths are used in the fit, the odd depths points are presented with x symbols simply to show their consistency with the model. Note: the green data is the same in both panels as we fix the incoherent errors to be zero.

3. Fitting CAFE data

In Fig. 13, we present examples of the CAFE data together with fits to the model of Eq. (4) realized on the simulation data to obtain the CZ error budgets presented in Figs. 11 and 12. We see that the model of Eq. (4) fits the CAFE data very well for realistic noise of strength $\epsilon_{\text{incoh}} = 0.005$ and $\epsilon_{\text{coh}} = 0.002$, even when the incoherent noise comes from an amplitude and phase damping channel. In Fig. 13, we also show simulation data using the same noise channel but an ideal unitary ($\epsilon_{\text{coh}} = 0$) in red, and using the same unitary without any incoherent noise ($\epsilon_{\text{incoh}} = 0$) in green. The linear and purely quadratic behaviors of these curves, respectively, can be easily understood from the quadratic form of Eq. (8) where the incoherent errors build up linearly with n , whereas the coherent errors build up quadratically.

Since we simulate the actual 16 circuits used to obtain the average gate fidelity at different depths, we capture some depth-dependent behaviors that are due to coherent effects in the single- and two-qubit unitaries (see for example the green ‘x’ at $n = 1$ in Fig. 13, which is actually higher than the point at $n = 0$). This is not an artifact of the finite sampling, but is due to the fact that the preparation and measurement circuits both require an imperfect CZ gate, which can coherently map the state closer or further away from the desired state $|00\rangle$, depending on the specific imperfect unitaries. Similarly, part of the incoherent noise channel can anti-commute with the CZ unitary and produce back-and-forth behaviors between the odd and even depths n . We have found in simulation that using only the even depths avoids this issue, while performing similarly or better in terms of accuracy in error budgeting, compared to using all the even and odd depths.

Appendix F: Dynamical Decoupling in CAFE

In this Appendix, we show how applying an X gate to both qubits after the CZ in the cycle circuit decouples the fSim unitary from both of its single-qubit phases. First, we can show that interleaving a pair of fSim gates with parallel X gates gives a cycle unitary that also corresponds to an fSim gate. To see this, consider what happens when conjugating an fSim unitary by parallel Xs:

$$(X \otimes X) \text{fSim}(X \otimes X) = \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & e^{-i\gamma} u_{11} & e^{-i\gamma} u_{12} & 0 \\ 0 & e^{-i\gamma} u_{21} & e^{-i\gamma} u_{22} & 0 \\ 0 & 0 & 0 & e^{-i(2\gamma+\phi)} \end{pmatrix} \begin{pmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \end{pmatrix} \quad (\text{F1})$$

$$= \begin{pmatrix} e^{-i(2\gamma+\phi)} & 0 & 0 & 0 \\ 0 & e^{-i\gamma} u_{22} & e^{-i\gamma} u_{21} & 0 \\ 0 & e^{-i\gamma} u_{12} & e^{-i\gamma} u_{11} & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix}, \quad (\text{F2})$$

where we’ve factored the fSim unitary into phases between the particle-number subspaces and a special unitary U in the single-excitation subspace

$$U = \begin{pmatrix} u_{11} & u_{12} \\ u_{21} & u_{22} \end{pmatrix} = \begin{pmatrix} e^{-i\zeta} \cos \theta & -ie^{i\chi} \sin \theta \\ -ie^{-i\chi} \sin \theta & e^{i\zeta} \cos \theta \end{pmatrix}. \quad (\text{F3})$$

Following this with a second fSim gate gives

$$\text{fSim}(X \otimes X) \text{fSim}(X \otimes X) = e^{-i2\gamma} \begin{pmatrix} e^{-i\phi} & 0 & 0 & 0 \\ 0 & \tilde{u}_{11} & \tilde{u}_{12} & 0 \\ 0 & \tilde{u}_{21} & \tilde{u}_{22} & 0 \\ 0 & 0 & 0 & e^{-i\phi} \end{pmatrix} \quad (\text{F4})$$

$$UXUX = \tilde{U} = \begin{pmatrix} \tilde{u}_{11} & \tilde{u}_{12} \\ \tilde{u}_{21} & \tilde{u}_{22} \end{pmatrix} = \begin{pmatrix} u_{11}u_{22} + u_{12}^2 & u_{11}(u_{21} + u_{12}) \\ u_{22}(u_{21} + u_{12}) & u_{11}u_{22} + u_{21}^2 \end{pmatrix}. \quad (\text{F5})$$

Here the common phase γ has been eliminated (only appearing as a global phase), and the controlled phase ϕ is a relative phase between the even- and odd-parity subspaces. For the case of something close to a CZ gate (small swap angle θ and small differential phase ζ), the single-excitation unitary simplifies to

$$U = I - i(\theta \cos \chi X - \theta \sin \chi Y + \zeta Z) + \mathcal{O}(\zeta^2) + \mathcal{O}(\theta^2) \quad (\text{F6})$$

$$XUX = I - i(\theta \cos \chi X + \theta \sin \chi Y - \zeta Z) + \mathcal{O}(\zeta^2) + \mathcal{O}(\theta^2) \quad (\text{F7})$$

$$UXUX = I - i2\theta \cos \chi X + \mathcal{O}(\zeta^2) + \mathcal{O}(\theta^2) + \mathcal{O}(\zeta\theta). \quad (\text{F8})$$

As such, the effects of the differential phase ζ , to first order, are also removed by this echoing.

Appendix G: Swap errors without phase matching

When running CAFE with characterized unitaries in the experiments performed in this work, we do not attempt to correct the swap errors. This is for two reasons. First, the swap errors are the smallest errors in the system, and tend to be overshadowed by the other sources of error. Second, in our system, the relative phases accumulated between qubits sitting at different idle frequencies are accounted for by changing the microwave phases. This is possible because the two-qubit gate we are using, the CZ gate, ideally commutes with such differential phases. The effect of not “phase matching” (that is, removing this differential phase by shifting qubit frequencies) on swap errors is that the swap phase χ is shifted from cycle to cycle, preventing the swaps from coherently adding and further diminishing their effect relative to the other error sources. Should one be concerned with relatively large swap angles θ , one could incorporate the application-dependent value of χ in the determination of the circuit to map the predicted state back to $|00\rangle$.

Appendix H: Additional CZ Data

For completeness, we present in Fig. 14 the entire error budget results from the parallel CZ CAFE dataset acquired on a Sycamore processor. This data illustrates that CAFE can be straightforwardly deployed on large-scale quantum processors to characterize the fidelity of quantum operations in parallel, while budgeting the incoherent and coherent error contributions. For reference, we also present an error budgeting of the same CZ gates obtained from XEB. Note that in this latter case, the coherent error contribution is obtained by subtracting the incoherent speckle purity error from the total XEB error, which gives an unphysical negative ϵ_{incoh} value for 33 out of the 103 gates. In our tests, this problem did not arise for CAFE.

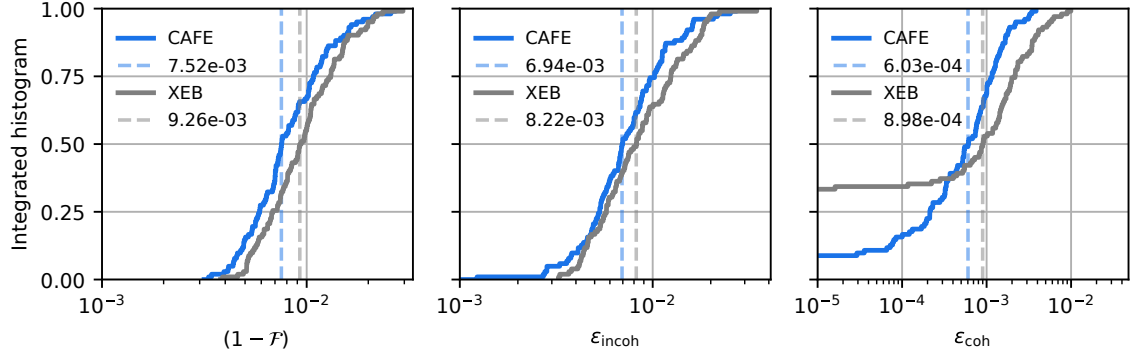


FIG. 14. Error budgets of parallel CZ gates on a Sycamore processor. Data in blue is obtained from the CAFE experiment using the procedure detailed in the main text, where 3 samples of this data are shown in Fig. 3. For reference, we present similar gate error budget information as obtained from standard XEB in gray, where the incoherent error is extracted from the speckle purity and the coherent error is obtained from the difference with the total gate error. Data shows 103 different two-qubit gates.

5.4 Publication [4] : Estimation des paramètres unitaires de portes quantiques (MEADD)

Je suis second auteur de l'article suivant [4], lequel est en processus de révision finale pour sa publication dans la revue *npj Quantum Information*. Alors que le premier et le dernier auteur ont eu l'idée du protocole MEADD et ont élaboré la majorité des formulations analytiques, j'ai principalement contribué à raffiner la méthodologie ainsi qu'à la mise en œuvre numérique et expérimentale du protocole pour l'ensemble de ses applications sur différentes portes quantiques. J'ai grandement contribué à rédiger l'article en plus d'acquérir toutes les données expérimentales et de produire l'ensemble des figures.

Characterizing Coherent Errors using Matrix-Element Amplification

Jonathan A. Gross,^{1,*} Élie Genois,^{1,2} Dripto M. Debroy,¹ Yaxing Zhang,¹ Wojciech Mruczkiewicz,¹ Ze-Pei Cui,^{1,3} and Zhang Jiang^{1,†}

¹*Google Quantum AI, Venice, CA 90291, USA*

²*Institut quantique & Département de Physique, Université de Sherbrooke, Quebec J1K2R1, Canada*

³*Department of Physics & Joint Quantum Institute, University of Maryland, MD 20742, USA*

(Dated: April 22, 2024)

Repeating a gate sequence multiple times amplifies systematic errors coherently, making it a useful tool for characterizing quantum gates. However, the precision of such an approach is limited by low-frequency noises, while its efficiency hindered by time-consuming scans required to match up the phases of the off-diagonal matrix elements being amplified. Here, we overcome both challenges by interleaving the gate of interest with dynamical decoupling sequences in a protocol we call Matrix-Element Amplification using Dynamical Decoupling (MEADD). Using frequency-tunable superconducting qubits from a Google Sycamore quantum processor, we experimentally demonstrate that MEADD surpasses the accuracy and precision of existing characterization protocols for estimating systematic errors in single- and two-qubit gates. In particular, MEADD yields factors of 5 to 10 improvements in estimating coherent parameters of the CZ gates compared to existing methods, reaching a precision below one milliradian. We also use it to characterize coherent crosstalk in the processor which was previously too small to detect reliably.

I. INTRODUCTION

Scalable quantum computation suffers from systematic errors, such as over-rotation and crosstalk. Characterizing these errors often requires coherently amplifying them by repeating the same gate sequence in a quantum circuit, a core principle of many schemes [1–5]. This approach also allows one to distinguish errors in the gate of interest from state preparation and measurement (SPAM) errors. However, in practice, noise in the system often overshadows small, systematic coherent errors, limiting the extent to which they can be amplified and characterized with precision. In particular, frequency-tunable superconducting qubits suffer from slow fluctuations and drift in qubit frequencies [6, 7], imposing a noise floor that limits the precision of standard characterization approaches.

In this work, we introduce Matrix-Element Amplification using Dynamical Decoupling (MEADD) – a protocol for amplifying coherent errors even in the presence of low-frequency noise. This is achieved by interleaving the gate of interest with carefully chosen single-qubit gates that implement a dynamical decoupling (DD) sequence [8–12], echoing away irrelevant and noisy parameters while amplifying the parameter being measured. The working principle of MEADD is illustrated in Fig. 1, using the CZ gate as an example. While its primary application is for characterizing two-qubit gates, we also apply the principles of MEADD to the single-qubit setting by incorporating the gate of interest into the DD sequence. We experimentally demonstrate the effectiveness of our approach using a Google Sycamore processor, see Supplementary Information in [13] for details of the hardware.

We combine ideas from several established characterization techniques. Like Floquet calibration [4, 5], MEADD fundamentally measures eigenvalue differences, allowing amplification of errors and mitigation of SPAM. The additional ingredient is to use DD sequences to ensure the eigenvalue difference only depends on one gate parameter at a time, much like what is done in Hamiltonian learning [14]. This isolation significantly simplifies data processing, making it easy to debug and relatively robust to unknown processes such as stray couplings and leakage. By echoing away unknown and irrelevant diagonal terms, it also eliminates the need for costly scans to identify resonance conditions, i.e., phase matching off-diagonal matrix elements so that they add to each other constructively when the gate is repeated n times. This gives MEADD greater efficiency in the number of experiments to run, which need only scale logarithmically in the maximum circuit depth, compared to similar noise-robust methods based on signal processing [15], which scale linearly in circuit depth.

An essential requirement of MEADD is the robustness to systematic errors in the DD implementation, such as over rotations. These errors can also be amplified coherently, ultimately spoiling the estimation of the intended parameter. We present a systematic approach to constructing DD sequences such that these types of errors cancel each other to the first order, ensuring accuracy and reliability in practice. Since these robust DD sequences depend on the gate of interest, MEADD is most suitable for characterizing gates whose errors can be thought of as perturbations of the ideal gate.

The main text begins with Sec. II, where we describe implementation details for single- and two-qubit gates on frequency-tunable superconducting qubits. Understanding these details is critical to designing efficient methods for characterizing such operations. In Sec. III, we present experimental results for the CZ gate, single-qubit X ro-

* E-mail: jarthurgross@google.com

† E-mail: qzj@google.com

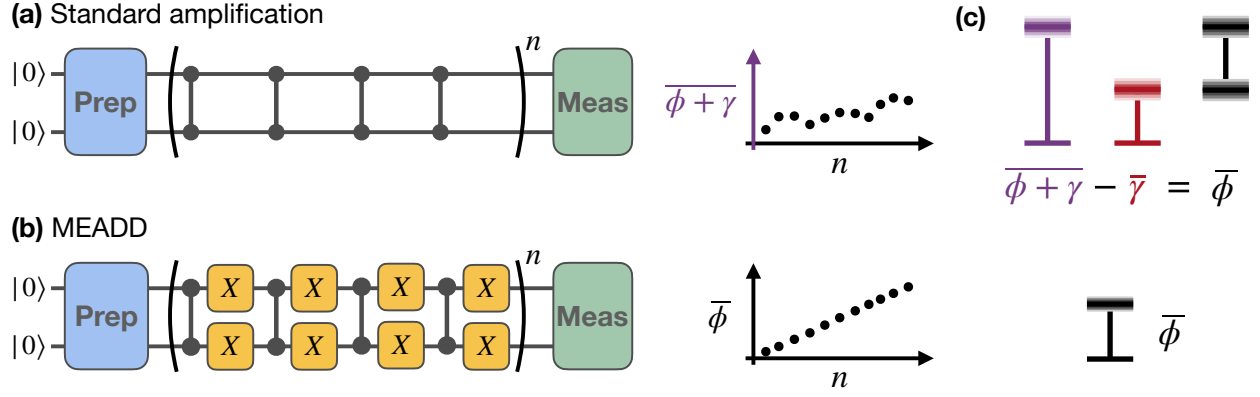


FIG. 1. Schematic of the working principle of MEADD. **(a)** Standard gate repetition amplifies the matrix elements to estimate, but also amplifies other coherent and incoherent (noise) processes. **(b)** Making use of single-qubit dynamical decoupling sequences, MEADD filters out specific matrix elements while additionally echoing out the noisy time-varying single-qubit phases. The parameter we are estimating—together with other matrix element commuting with the DD gates—then builds up linearly with depth and can be fitted from the slope, with Heisenberg scaling in precision with circuit depth. All MEADD circuits possess this simple structure, with a preparation and pre-measurement layer containing at most a single entangling operation together with at most four single-qubit gates, and terminated by standard single-qubit Z-basis measurement. **(c)** Fluctuations in a noisy parameter γ (red), representing single-qubit phases for example, can overshadow the estimation of a stable parameter ϕ (black) in a simple gate repetition scheme that provides estimates depending on both of these parameters (purple). Note that a separate, and necessarily noisy, estimation of γ is therefore required to find ϕ . In contrast, MEADD provides a direct estimate of the stable parameter ϕ as illustrate on the bottom, thus yielding more precise and accurate characterization information.

tations, and crosstalk. In all these cases, MEADD gives improved accuracy, precision, and efficiency over existing techniques. In Sec. IV we discuss the characterization protocol for single-qubit gates and explain how matrix-element amplification works in practice. In Sec. V we detail the MEADD protocol for characterizing the entangling parameters of two-qubit excitation-preserving gates before demonstrating in Sec. VI that it is robust to implementation errors in the DD gates.

II. GATES IN SUPERCONDUCTING SYSTEMS

Leveraging the hierarchy of precision in the different controls available on a quantum device is central to designing highly accurate and reliable characterization protocols. It is important to understand how quantum gates are implemented with specific types of devices. To that end, we present the considerations for calibrating gates in tunable-frequency superconducting qubits, highlighting the principal noise sources therein [7, 16–18]. That said, our method can be adapted to various types of quantum devices by modifying the gate sequences accordingly.

Consider a qubit driven by a microwave pulse, where we ignore leakage errors into energy levels lying outside of the qubit subspace. Under the rotating wave approximation (RWA), the Hamiltonian for this two-level system is ($\hbar = 1$)

$$H_{\text{RWA}}(t) = -\frac{1}{2}\omega_{01}(t)Z + \left(\Omega(t)e^{-i\nu t}\sigma_- + \text{h.c.}\right), \quad (1)$$

where $Z = |0\rangle\langle 0| - |1\rangle\langle 1|$, $|0\rangle$ and $|1\rangle$ representing the

ground and first excited states of the qubit, respectively. The terms with $\sigma_- = |1\rangle\langle 0|$ and its Hermitian conjugate $\sigma_+ = |0\rangle\langle 1|$ capture the microwave control field, where $\Omega(t)$ is an envelope function with bandwidth ~ 100 MHz and $\nu/2\pi$ is the nominal microwave frequency, sitting at several GHz. The parameter $\omega_{01}(t)$ is the angular frequency of the qubit, adjustable via the SQUID loop flux. Ideally, $\omega_{01}(t)$ matches ν during idle periods and rotations about axes in the X-Y plane. However, noise in the flux control line and surrounding electromagnetic environment can cause $\omega_{01}(t)$ to drift, leading to various errors when performing qubit operations. This can be seen directly by moving to the interaction frame of the microwave frequency $|\psi_I(t)\rangle = e^{-i\nu t Z/2} |\psi(t)\rangle$, which eliminates the rapidly oscillating phase $e^{-i\nu t}$ in Eq. (1) as follows

$$H_I(t) = e^{-i\nu t Z/2} H_{\text{RWA}}(t) e^{i\nu t Z/2} + \frac{1}{2}\nu Z \quad (2)$$

$$= -\frac{1}{2}\delta\omega_{01}(t)Z + \left(\Omega(t)\sigma_- + \text{h.c.}\right). \quad (3)$$

Here, one can directly see how any non-zero detuning $\delta\omega_{01}(t) = \omega_{01}(t) - \nu \ll \nu$ can lead to unwanted Z rotations. The resulting unitary produced by this drive can be formally written as

$$U_I = \mathcal{T} \exp\left(-i \int_{t_i}^{t_f} H_I(t) dt\right), \quad (4)$$

where $\mathcal{T}$ is the time ordering operator and t_i (t_f) the initial (final) time of the pulse.

We now discuss how to implement quantum logical circuits in the above interaction picture, i.e. control $H_I(t)$ such that U_I corresponds to the desired quantum gate. Note that as we always initialize and measure the qubits in the Z basis, there is no need to implement the unitary $e^{-i\nu t Z/2}$ for the frame change before the final measurement. However, any gate operation that is not diagonal in the Z basis needs to account for the frame change, a process known as phase tracking or matching.

Single-qubit XY rotations (microwave gates). An arbitrary SU(2) matrix can be parameterized using three angles as

$$R(\mu, \zeta, \chi) = \begin{pmatrix} e^{-i\zeta} \cos \frac{\mu}{2} & -i e^{i\chi} \sin \frac{\mu}{2} \\ -i e^{-i\chi} \sin \frac{\mu}{2} & e^{i\zeta} \cos \frac{\mu}{2} \end{pmatrix}. \quad (5)$$

This parametrization directly leads to an Euler decomposition of arbitrary single-qubit gates as an X rotation sandwiched by two Z rotations:

$$R(\mu, \zeta, \chi) = Z(\zeta - \chi) X(\mu) Z(\zeta + \chi), \quad (6)$$

where

$$X(\varphi) = e^{-i\varphi X/2}, \quad Z(\varphi) = e^{-i\varphi Z/2}. \quad (7)$$

The rotation $X(\mu)$ can be achieved by applying a microwave pulse to the qubit, e.g., $\Omega(t) = \mu C(t)$, where $C(t)$ is the shifted cosine function

$$C(t) = \begin{cases} \frac{1}{2T} (1 + \cos \frac{2\pi t}{T}), & \text{for } -\frac{T}{2} \leq t \leq \frac{T}{2}, \\ 0, & \text{for } |t| > \frac{T}{2}, \end{cases} \quad (8)$$

where the pulse time T is around 15 ns in our system. As such, the parameter μ of the single-qubit gate is directly determined by the amplitude of the microwave drive, which can be controlled to a high precision using standard arbitrary waveform generators.

More generally, we consider a rotation along an arbitrary axis in the XY plane, which we notate as a *phased-X* rotation

$$\text{PhX}(\mu, \vartheta) \equiv Z(\vartheta) X(\mu) Z(-\vartheta), \quad (9)$$

where μ and ϑ specify the angle and axis of the rotation, respectively. It can be implemented by simply changing the phase of the microwave pulse

$$\Omega(t) = \mu e^{i\vartheta} C(t). \quad (10)$$

This microwave phase can be used to control χ much more precisely than the uncertainties on the qubit phase, making it an ideal control knob to use in designing characterization protocols. Additionally, the freedom we have in the microwave phase allows us to eliminate one Z rotation from our single-qubit gate

$$R(\mu, \zeta, \chi) = Z(2\zeta) \text{PhX}(\mu, -\zeta - \chi), \quad (11)$$

which can save time and reduce errors.

In practice, we additionally make corrections to the pulse, such as compensating for the AC Stark effect by introducing a detuning $\Delta(\mu)$,

$$\Omega(t) = \mu e^{i\vartheta} e^{-i\Delta(\mu)t} C(t), \quad (12)$$

which can also be used to adjust the gate parameter ζ . Finally, we make use of Derivative Removal by Adiabatic Gate (DRAG) [19] pulses to mitigate leakage errors, which consist of mixing the original pulse $\Omega(t)$ with its first time derivative; see App. A for a simple derivation. Impedance mismatches in the control lines can reflect microwave pulses, which cause delayed tails after the main pulses. Here we assume that these tails do not bleed into the entangling pulses such as to preserve the independence of consecutive gates.

Physical single-qubit Z rotations. Rotations around the Z axis can be implemented simply by moving the qubit frequency away from ν for some duration using a flux pulse, while leaving $\Omega(t) = 0$. This can be seen directly by considering the evolved qubit state in the interaction frame under the propagator of Eq. (4)

$$|\psi_I(t_f)\rangle = e^{i\varphi(t_f, t_i)Z/2} |\psi_I(t_i)\rangle, \quad (13)$$

where the Z-rotation angle reads

$$\varphi(t_f, t_i) = \int_{t_i}^{t_f} \delta\omega_{01}(t) dt. \quad (14)$$

This type of Z rotation typically takes ~ 20 ns to implement on our hardware (including a ~ 10 ns pulse and two 5 ns paddings before and after the pulse), regardless of the angle, with a corresponding gate infidelity of about 10^{-3} . However, these physical Z pulses typically have nonzero tails lasting much longer than the gate time [20–22], on the order of hundreds of ns in our devices, which can disturb many subsequent operations.

Virtual single-qubit Z rotations. For quantum circuits consisting of only single-qubit gates and diagonal two-qubit gates, such as the CZ gate, the Z rotations can be implemented in software without actually sending flux pulses to the qubits [17]. This is because Z rotations commute with diagonal two-qubit gates, and when commuted through a microwave gate, only change its phase χ , which we can control with high precision relative to physical Z rotations.

In more detail, one can push single-qubit Z rotations from the beginning of the circuit down—through any two-qubit gates—and combine them with other Z rotations along the way until they hit single-qubit gates that are not Z rotations. One can then use Eq. (11) to put the Z rotations after the XY rotations, and repeat the same procedure until all the single-qubit Z rotations are pushed to the end of the circuit. Finally, they can be directly eliminated because the qubits are measured in the Z basis. The above process also works for two-qubit gates that are diagonal up to a SWAP. However, physical Z rotations are still required for most other two-qubit

gates, and they are typically applied before every two-qubit pulses in these cases.

Relative axes of XY rotations. Consider a circuit consisting of only XY rotations and diagonal two-qubit gates, such as CZ. One can implement the following gauge transformation on all the XY rotations without affecting the measurement results

$$\chi_j \rightarrow \chi_j + \text{constant}, \quad (15)$$

where χ_j is the rotation axis of the j -th XY gate. Therefore, the absolute value of χ has no physical significance and cannot be measured using such circuits. The relative χ between two different XY rotations, however, is physical and needs to be calibrated. For example, we consider the relative χ between a π gate and a $\pi/2$ gate, which is nonzero due to effects such as ac-Stark shifts imparting χ differently for different values of μ . Suppose one has a well calibrated $-\pi/2$ pulse, meaning $\mu = -\pi/2$ and ζ has been removed through virtual Z phase updates,

$$R(-\pi/2, 0, \chi_{\pi/2}) = Z(-\chi_{\pi/2}) X(-\pi/2) Z(\chi_{\pi/2}), \quad (16)$$

and a π pulse with well calibrated rotation angle μ

$$R(\pi, 0, \chi_\pi) = Z(-\chi_\pi) X(\pi) Z(\chi_\pi). \quad (17)$$

One can characterize the relative phase $\chi_{\pi/2} - \chi_\pi$ by considering the composite gate:

$$R(\pi, 0, \chi_\pi) R(-\pi/2, 0, \chi_{\pi/2}) \quad (18)$$

$$= Z(\chi_{\pi/2} - 2\chi_\pi) X(\pi/2) Z(\chi_{\pi/2}) \quad (19)$$

$$= R(\pi/2, \zeta, \chi), \quad (20)$$

where

$$\zeta = \chi_{\pi/2} - \chi_\pi, \quad \chi = \chi_\pi. \quad (21)$$

The relative phase is now ζ of the composite gate, and can be estimated using our calibration circuits.

Two-qubit gates. In a Sycamore quantum processor, two-qubit gates are performed by applying flux pulses to both qubits and the frequency-tunable coupler between them [23]. The first two pulses bring the qubits into resonance with each other, while the third pulse activates a tunable excitation-preserving interaction that can be used to perform entangling gates. We parametrize the most general excitation-preserving two-qubit gates by the swap angle $0 \leq \theta \leq \pi/2$, the swap phase χ , the controlled phase ϕ , and the differential and common single-qubit phases ζ and γ , respectively:

$$U(\theta, \zeta, \chi, \gamma, \phi) = \begin{pmatrix} e^{i\gamma} & 0 & 0 & 0 \\ 0 & e^{-i\zeta} \cos \theta & -i e^{i\chi} \sin \theta & 0 \\ 0 & -i e^{-i\chi} \sin \theta & e^{i\zeta} \cos \theta & 0 \\ 0 & 0 & 0 & e^{-i(\gamma+\phi)} \end{pmatrix}. \quad (22)$$

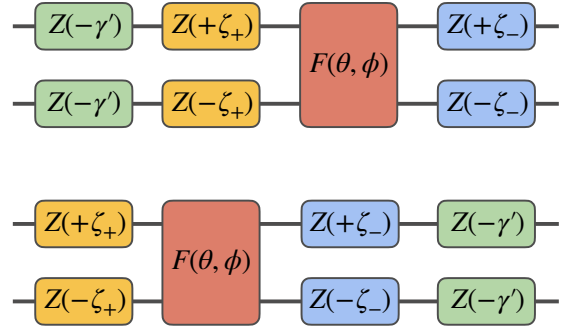


FIG. 2. Two equivalent ways to decompose the arbitrary excitation-preserving two-qubit gate $U(\theta, \zeta, \chi, \gamma, \phi)$ defined in Eq. (22) into single-qubit Z rotations and the fundamental entangler $F(\theta, \phi) = e^{-i\theta(X \otimes X + Y \otimes Y)/2 - i\phi Z \otimes Z/4}$. While the differential phases $\zeta_{\pm} = (\zeta \pm \chi)/2$ do not commute with $F(\theta, \phi)$, the symmetric phase $\gamma' = \gamma + \phi/2$ does and can be placed either before or after the entangling operation.

In our representations of two-qubit gates, we use the lexicographically ordered basis $\{|00\rangle, |01\rangle, |10\rangle, |11\rangle\}$, where the left bit corresponds to the left qubit in tensor products and the top wire in circuit diagrams.

One can separate the entangling parameters θ and ϕ from the single-qubit Z-phase parameters χ , ζ , and γ , as show in Fig. 2, which is a special case of the KAK decomposition [24]. Given the aforementioned low-frequency flux noise and drift in the control electronics, the parameters ζ , χ , and γ associated with single-qubit Z phases are much less stable than the entangling parameters θ and ϕ in our device. The instability in single-qubit phases ultimately hinders one's ability to characterize precisely the entangling-gate parameters, as shown schematically in Fig. 1(c). We solve this problem with MEADD by decoupling the estimation of θ and ϕ from the noise associated with flux control.

III. EXPERIMENTAL RESULTS

In this section we present comprehensive benchmark results on a 69-qubit Sycamore processor for characterizing CZ gates, single-qubit microwave gates, and crosstalk due to charge-charge (capacitive) coupling. In all these cases MEADD significantly outperforms other characterization methods both in terms of accuracy, precision, and efficiency. We defer the implementation details and robustness discussions to later sections.

Entangling parameters in CZ. We compare MEADD to cross-entropy benchmarking (XEB) [13], unitary tomography (UT) [23] and the phase method (PM) [5] in Fig. 3. In panel (a), we use a recently introduced protocol named Context-Aware Fidelity Estimation (CAFE) [25] to assess the accuracy of different methods, where the average gate fidelity between the learned unitary and the physical gate implemented is estimated. Using different numbers of gate repetitions n gives the protocol robustness

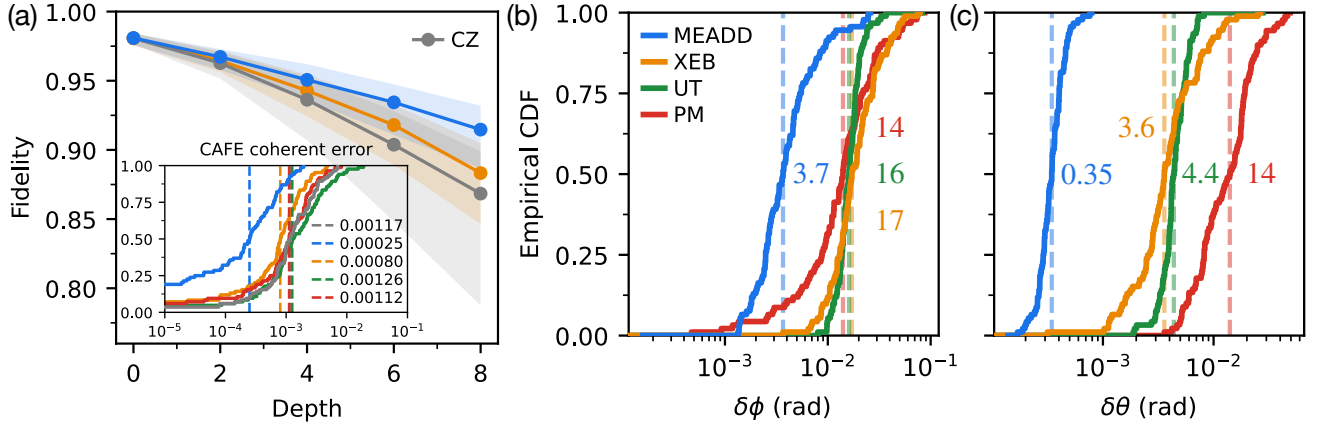


FIG. 3. Experimental results of MEADD on CZ gates. **(a)** Fidelity of our CZ gates obtained using the Context-Aware Fidelity Estimation (CAFE) protocol [25] with respect to unitaries characterized by MEADD (blue), Cross-Entropy Benchmarking (XEB, orange), and a perfect CZ (gray). Solid lines with markers are the median over 85 different CZ gates of a single Sycamore device, and the shaded regions are the inner quartile regions. A higher CAFE fidelity corresponds to a more accurate unitary characterization of the gate, see text for details. Inset shows the empirical cumulative distribution functions (CDF) of the coherent errors inferred from the CAFE curves of every CZ gate, with median values reported in the legend. Also shown are the results obtained using the unitary tomography (UT, green) and the phase method (PM, red) characterization protocols. **(b)** Empirical CDFs of run-to-run standard deviations in the characterized controlled phase ϕ , and **(c)** swap angle θ , measured over 12 sequential characterizations of the 85 CZ gates. The different protocols were interleaved with one another over several hours to average out device inconsistencies. The median standard deviations of the different methods are reported in milliradian in the plots.

to SPAM errors and allows for error budgeting in terms of incoherent and coherent contributions, as these errors build up linearly and quadratically with n to lowest order [25]. As such, the CAFE results at each depth n give an estimate of the average gate fidelity between any given unitary U^n and the physical gates implemented (nominally CZ^n) such that a higher fidelity directly corresponds to a more accurate unitary description of the gate. The unitaries extracted with MEADD outperform all other unitary characterization techniques as measured by CAFE. Specifically, looking at the inset of Fig. 3(a), we see that MEADD captures significantly more of the coherent error budget, resulting in a median of 2.5×10^{-4} uncharacterized coherent error across the 85 CZ gates characterized, 3-5x smaller than the other methods.

In Fig. 3(b-c), we compare the precision of MEADD to the other methods by repeating the different unitary characterization techniques and reporting the run-to-run standard deviation of the estimated two-qubit entangling parameters ϕ and θ for all 85 CZ gates of the device. MEADD provides significantly more consistent characterizations than any other method, with 5x and 10x improvements over XEB for the controlled phase ϕ and swap angle θ , respectively. Remarkably, MEADD achieves a 3.4×10^{-4} rad standard deviation for estimating the residual swap angle θ of our CZ gates. Such precision, below one milliradian, promises to be a key asset in the calibration of stable high fidelity two-qubit gates as required for scalable quantum error correction applications. Combined with the high accuracy of the MEADD estimates shown in Fig. 3(a), these results show the effectiveness

of dynamical decoupling pulses in isolating the parameter of interest from low-frequency and other noise sources present in frequency-tunable transmons.

Relative axes of XY rotations. We additionally show how low-frequency noise is mitigated using a simple example that can be regarded as MEADD for single-qubit gates. In Fig. 4, we present similar benchmarks of the accuracy and precision of MEADD, now for the case of interleaved single-qubit $X(-\pi/2)$ and $X(\pi)$ gates over the 69 qubits of the device. In panel (a), we compare the CAFE results obtained using the unitaries constructed from the MEADD estimates (blue) of the relative axis $\zeta = \chi_{\pi/2} - \chi_{\pi}$ in Eq. (21) to the ones assuming perfect gates (gray). The higher fidelities obtained from the MEADD unitaries indicate that these characterizations accurately describe the implemented gates. In panels (b) and (c), we compare the MEADD characterization results to a direct 4-axis tomography protocol which consists of preparing an equatorial state, applying an $X(\pi)$ pulse, and then measuring expectation values of X and Y . Since the state preparation and measurements use $X(-\pi/2)$ gates, this protocol measures the relative axes of the $X(\pi)$ and $X(-\pi/2)$ gates just like MEADD with interleaved $X(-\pi/2)$ and $X(\pi)$ gates does, as described at the end of Sec. IV. MEADD achieves a median run-to-run standard deviation of 6.11×10^{-4} rad on estimating ζ , a 6x improvement over the tomography protocol. This is illustrated in panel (b) for a representative qubit, showing the tight distribution of repeated characterizations. By providing such reliable unitary characterizations in an experimentally efficient protocol, MEADD demonstrates

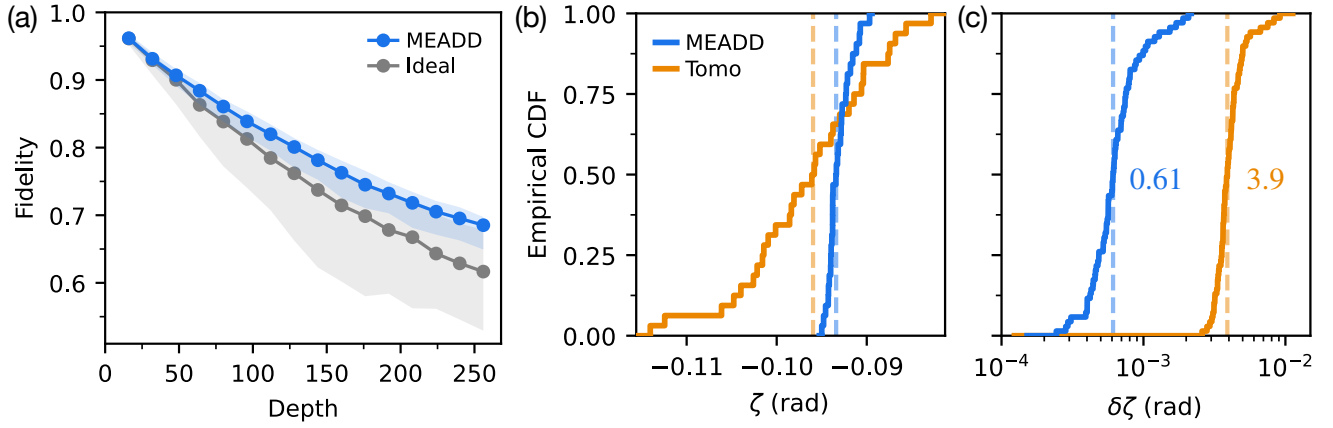


FIG. 4. Experimental results of MEADD for interleaved single-qubit $X(-\pi/2)$ and $X(\pi)$ gates. (a) CAFE results for the single-qubit unitaries characterized by MEADD (blue) compared to assuming a perfect gate (gray) for 69 qubits of a single Sycamore device. The cycle circuit being repeated in CAFE is the same as in the MEADD protocol, namely a $X(-\pi/2)$ rotation followed by a $X(\pi)$ rotation, performed in parallel on all qubits. Solid lines with markers are the median over the 69 qubits, and the shaded regions are the inner quartile regions. (b) Distributions of the characterized relative rotation axis parameter $\zeta = \chi_{\pi/2} - \chi_{\pi}$ defined in Eq. (21) for a representative single-qubit gate characterized 32 times interleaved between MEADD (blue) and 4-axis tomography (orange), using the same number of measurements for both methods. (c) Distributions of run-to-run standard deviations in ζ estimates for the two methods over the 32 sequential characterizations of the 69 qubits. The median standard deviations of the two methods are reported in milliradian.

its utility for calibrating high-fidelity operations in qubits subject to low-frequency flux noise.

Crosstalk. While we have discussed MEADD so far in the context of characterizing gates, we can also use it to measure other parameters of our system. A good example is crosstalk between qubits caused by unwanted charge-charge (capacitive) couplings, which results in a qubit-subspace Hamiltonian of the form

$$H \propto e^{i\chi} \sigma_+ \otimes \sigma_- + \text{h.c.}, \quad (23)$$

and produces an unwanted excitation swapping with angle θ_{xtalk} . These terms are often hard to measure, as swapping between the two qubits is suppressed when their detuning, on the order of 100 MHz, is much larger than the unwanted coupling, which is typically on the order of 100 kHz. The errors this incurs for typical gate times (10s of nanoseconds) yields average gate infidelities between 1×10^{-5} and 1×10^{-4} , which while small can be particularly damaging due to their nonlocal nature within error-correcting circuits [26].

A typical strategy for measuring the swap in this case is to tune the qubits to resonance. Since the qubit frequency uncertainty is typically on the order of 100 kHz, such a strategy requires scanning over the qubit-frequency control to find the resonance, which is expensive to perform in practice. In contrast, we can use MEADD to directly estimate the unwanted crosstalk without any qubit flux tuning.

We demonstrate the precision of this MEADD-based crosstalk characterization strategy in Fig. 5. As shown in panel (a), we obtain MEADD data with an unambiguous signal for swapping angles on the order of only 1 mrad. We can fit this data to reliably characterize crosstalk-

induced swapping with a precision below 0.1 mrad, as shown in panel (c), which corresponds to an uncertainty on the effective swapping coupling strength of less than 0.5 MHz. Note that this protocol can readily be used to characterize long-range crosstalk across distant qubits on the processor.

IV. CHARACTERIZING SINGLE-QUBIT GATES WITH COHERENT AMPLIFICATION

Having presented the experimental results of MEADD, we now describe how the protocol works, starting with the single-qubit version. For the arbitrary single-qubit gate $R(\mu, \zeta, \chi)$ in Eq. (5) we set $\mu = 2\theta$ to draw an explicit connection to the two-qubit parametrization (22),

$$R(\mu, \zeta, \chi) = I \cos \Omega(\theta, \zeta) - i \sigma(\theta, \zeta, \chi) \sin \Omega(\theta, \zeta), \quad (24)$$

where

$$\cos \Omega = \cos \theta \cos \zeta, \quad \Omega \in [\theta, \pi - \theta], \quad (25)$$

and the spin-1/2 operator

$$\sigma = \frac{\sin \theta}{\sin \Omega} (X \cos \chi - Y \sin \chi + Z \cot \theta \sin \zeta). \quad (26)$$

In matrix form, we have

$$\sigma(\theta, \zeta, \chi) = \begin{pmatrix} \cos \varrho & e^{i\chi} \sin \varrho \\ e^{-i\chi} \sin \varrho & -\cos \varrho \end{pmatrix}, \quad (27)$$

where the polar angle ϱ is defined by

$$\varrho = \text{arccot}(\cot \theta \sin \zeta), \quad \varrho \in [\theta, \pi - \theta]. \quad (28)$$

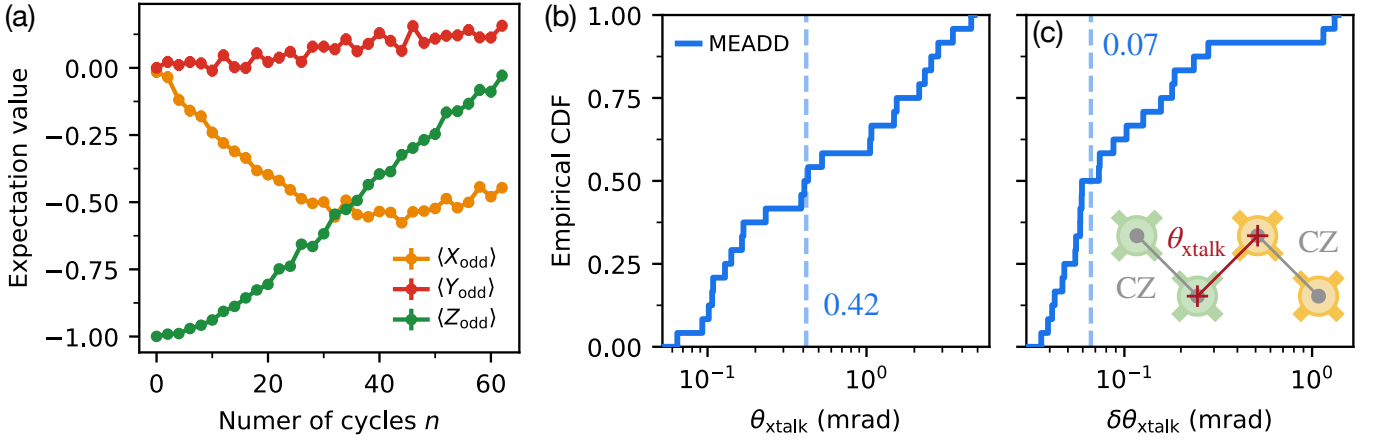


FIG. 5. Experimental results of MEADD for characterizing crosstalk. **(a)** Evolution of the components of the Bloch vector in the odd-parity sector, $X_{\text{odd}} = |10\rangle\langle 01| + |01\rangle\langle 10|$, $Y_{\text{odd}} = i|10\rangle\langle 01| - i|01\rangle\langle 10|$, and $Z_{\text{odd}} = |01\rangle\langle 01| - |10\rangle\langle 10|$, for two neighboring qubits involved in $2n$ parallel CZ gates. The corresponding crosstalk angle is about 4 mrad per cycle. Error bars correspond to the standard deviation over 2000 shots and are smaller than the data markers. Distributions of **(b)** the median crosstalk swap angle θ_{xtalk} and **(c)** the run-to-run standard deviation in this estimate after repeating the measurement 10 times over one hour for 24 neighboring CZ gates on the Sycamore processor. As shown in the inset, the 34 ns CZ gates (gray) are performed simultaneously on neighboring qubit pairs (green and yellow), and we use MEADD to characterize the residual swap angle produced by these operations, yielding $\theta_{\text{xtalk}} = (0.42 \pm 0.07)$ mrad on median.

We can measure the eigenvalues $\pm\Omega(\theta, \zeta)$ of R precisely by preparing and measuring uniform superpositions of the eigenstates of R :

$$|\sigma^+\rangle = e^{i\chi/2} \cos(\varrho/2) |0\rangle + e^{-i\chi/2} \sin(\varrho/2) |1\rangle \quad (29)$$

$$|\sigma^-\rangle = e^{i\chi/2} \sin(\varrho/2) |0\rangle - e^{-i\chi/2} \cos(\varrho/2) |1\rangle, \quad (30)$$

where $\sigma |\sigma^\pm\rangle = \pm |\sigma^\pm\rangle$. We parametrize these superpositions by the relative phase τ between them

$$|\Phi(\tau)\rangle = \frac{1}{\sqrt{2}} (e^{-i\tau/2} |\sigma^+\rangle + e^{i\tau/2} |\sigma^-\rangle). \quad (31)$$

After applying the cycle unitary n times to the initial state, we have

$$R^n |\Phi(\tau_0)\rangle = |\Phi(\tau_0 + 2n\Omega)\rangle. \quad (32)$$

We then measure the probability of finding $|\Phi(0)\rangle$,

$$P_0 = |\langle \Phi(0) | \Phi(\tau_0 + 2n\Omega) \rangle|^2 \quad (33)$$

$$= \frac{1 + \cos(\tau_0 + 2n\Omega)}{2}, \quad (34)$$

and the probability to find $|\Phi(\pi/2)\rangle$ is

$$P_{\pi/2} = |\langle \Phi(\pi/2) | \Phi(\tau_0 + 2n\Omega) \rangle|^2 \quad (35)$$

$$= \frac{1 + \sin(\tau_0 + 2n\Omega)}{2}. \quad (36)$$

Estimating these two measurement probabilities for a sequence of repetitions will give us the measurement record

$$2P_0 - 1 + i(2P_{\pi/2} - 1) = \exp[i(\tau_0 + 2n\Omega)], \quad (37)$$

from which we can extract Ω by fitting the complex argument as a linear function of repetitions. This procedure is robust to SPAM errors.

To extract the gate parameters μ and ζ , we perform a set of similar experiments, where we add different virtual Z rotations between the gates, giving us unitaries

$$Z(z)R(\mu, \zeta, \chi) = Z(\zeta + z - \chi) X(\mu) Z(\zeta + \chi) \quad (38)$$

$$= R(\mu, \zeta + z/2, \chi - z/2). \quad (39)$$

By measuring $\Omega(\theta, \zeta + z/2)$ as a function of z , one can fit the curve to the form given in Eq. (25), obtaining μ and ζ . Note this technique is identical to the technique for extracting two-qubit gate parameters within the single-excitation subspace given in [4]. Like there, the angle χ does not manifest in the eigenphase $\Omega(\theta, \zeta)$, and so we cannot coherently amplify errors in that parameter without additional interventions. In other words, χ corresponds to a gauge transformation

$$R(\mu, \zeta, \chi)^n = Z(-\chi) R(\mu, \zeta, 0)^n Z(\chi), \quad (40)$$

which does not add up when the same gate is repeated. That being said, the relative axis of two different single-qubit gates can be estimated with amplified precision. For example, we amplify the relative axis between $X(\pi)$ and $X(-\pi/2)$ defined in Eq. (21) by interleaving them many times. Any low-frequency noise in the qubit frequency that would cause phase to accumulate between the gates will be echoed away, making the measurement of this relative phase robust to these fluctuations.

V. CHARACTERIZING EXCITATION-PRESERVING TWO-QUBIT GATES WITH MEADD

The entangling parameters ϕ , θ , and χ of our two-qubit gates are more consistent over time than the single-qubit phases that are susceptible to low-frequency flux noise. It would therefore be useful to design characterization circuits that are sensitive only to the stable entangling parameters. To that end, we introduce microwave gates that are interleaved with the repeated two-qubit gate, as illustrated in Fig. 1. To remove flux noise, MEADD is designed to be insensitive to the single-qubit parameters ζ and γ ; we use an adapted Floquet characterization technique without DD to estimate these parameters, as described in Appendix B.

For the simple case of interleaving a two-qubit unitary U with X gates, one obtains the cycle unitary

$$C = (X \otimes X)U. \quad (41)$$

A useful representation for the following discussion is given by decomposing the two-qubit unitary from Eq. (22) into the even and odd parity subspaces $\mathcal{H}_{\text{even}} \oplus \mathcal{H}_{\text{odd}}$,

$$U = (e^{-i\phi/2} U_{\text{even}}) \oplus U_{\text{odd}} \quad (42)$$

$$= \left[e^{-i\phi/2} \begin{pmatrix} e^{i(\gamma+\phi/2)} & 0 \\ 0 & e^{-i(\gamma+\phi/2)} \end{pmatrix} \right] \oplus \begin{pmatrix} e^{-i\zeta} \cos \theta & -i e^{i\chi} \sin \theta \\ -i e^{-i\chi} \sin \theta & e^{i\zeta} \cos \theta \end{pmatrix}. \quad (43)$$

This decomposition is a direct sum of two $\text{SU}(2)$ unitaries with a relative phase $e^{-i(\phi/2)}$ between the even and odd parity sectors. Conveniently, $X \otimes X$ decomposes into the direct sum $X \oplus X$ in this representation, giving the decomposition for the cycle unitary

$$C = (e^{-i\phi/2} C_{\text{even}}) \oplus C_{\text{odd}}, \quad (44)$$

where $C_{\text{even}} = XU_{\text{even}}$ and $C_{\text{odd}} = XU_{\text{odd}}$. Importantly, the 2-cycle unitary is trivial in the even-parity subspace

$$C_{\text{even}}^2 = XU_{\text{even}}XU_{\text{even}} = I, \quad (45)$$

a feature we use to design our characterization protocol.

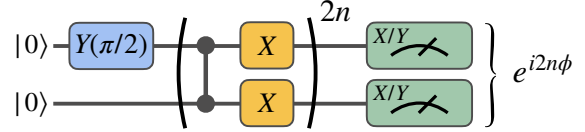
Controlled phase. For measuring ϕ , we need to measure a relative phase between the even- and odd-parity subspaces. Because the 2-cycle unitary is trivial on the even-parity subspace, we can choose any state such as $|00\rangle$ to serve as our even-parity reference.

To measure the relative phase $-\phi/2$ we will therefore prepare states of the form

$$|\psi_0\rangle = \frac{1}{\sqrt{2}}(|00\rangle + |\psi_{0,\text{odd}}\rangle) \quad (46)$$

$$= \frac{1}{\sqrt{2}}(|00\rangle + \alpha|01\rangle + \beta|10\rangle), \quad (47)$$

(a) ϕ



(b) θ, χ

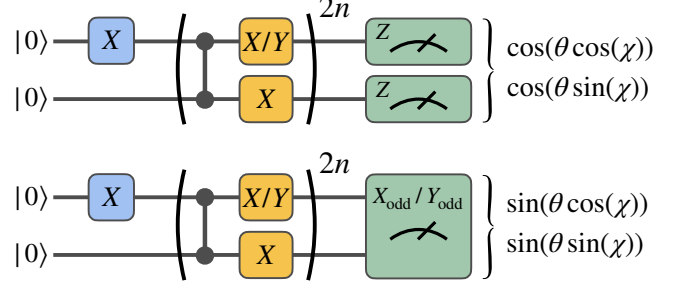


FIG. 6. MEADD circuits used to characterize the entangling parameters of CZ gates. (a) The circuits are composed of three steps. Prepare $|+0\rangle$ with a microwave gate (blue), or similarly $|0+\rangle$ (not shown), interleave $X \otimes X$ dynamical decoupling gates (yellow) with the repeated CZ gates for even repetitions $2n$, and finally measure both qubits in the X and Y bases (green). As explained in the main text, combining the resulting measurement bits for different depths n yields an estimate of the controlled phase ϕ that is robust to both low-frequency noise and DD gate imperfections. (b) Similarly for the swapping angle θ and phase χ , we prepare $|10\rangle$ (or $|01\rangle$, not shown), use $X \otimes X$ and $Y \otimes X$ as DD gates, and measure either in the computational basis of both qubits (top circuit) or in the Bell basis of the odd parity subspace (bottom circuit), where $X_{\text{odd}} = |10\rangle\langle 01| + |01\rangle\langle 10|$ and $Y_{\text{odd}} = i|10\rangle\langle 01| - i|01\rangle\langle 10|$.

where $|\psi_{0,\text{odd}}\rangle \in \mathcal{H}_{\text{odd}}$. Applying even powers of the cycle unitary to this initial state puts a phase on $|00\rangle$ and transforms the odd-parity part of the state separately

$$C^{2n}|\psi_0\rangle = \frac{1}{\sqrt{2}}(e^{-in\phi}|00\rangle + \alpha'|01\rangle + \beta'|10\rangle), \quad (48)$$

where

$$\alpha' = \alpha\langle 01|C_{\text{odd}}^{2n}|01\rangle + \beta\langle 01|C_{\text{odd}}^{2n}|10\rangle, \quad (49)$$

$$\beta' = \alpha\langle 10|C_{\text{odd}}^{2n}|01\rangle + \beta\langle 10|C_{\text{odd}}^{2n}|10\rangle. \quad (50)$$

We will make separable measurements on the qubits to infer ϕ , which one can understand by looking at the final density matrix $C^{2n}|\psi_0\rangle\langle\psi_0|C^{2n\dagger}$ and taking partial traces. Tracing out the right qubit gives

$$\begin{pmatrix} 1 + |\alpha'|^2 & e^{-in\phi}\beta'^* \\ e^{in\phi}\beta' & |\beta'|^2 \end{pmatrix}, \quad (51)$$

and tracing out the left qubit gives

$$\begin{pmatrix} 1 + |\beta'|^2 & e^{-in\phi}\alpha'^* \\ e^{in\phi}\alpha' & |\alpha'|^2 \end{pmatrix}. \quad (52)$$

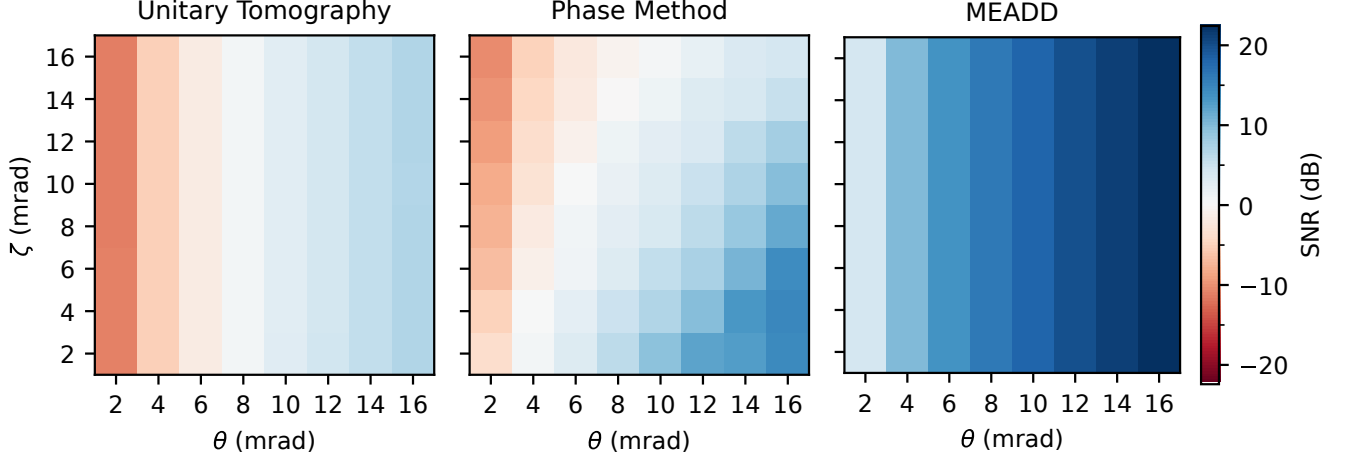


FIG. 7. Numerical simulation results on signal to noise ratio (SNR) for measuring the swap angle. Variance is measured over 100 different realizations for each parameter point. The number of simulated measurements is normalized so that the total number of measurements per parameter inferred is the same across the methods (see Appendix C for details). The phase method [5] has an advantage over unitary tomography in that it coherently amplifies errors. However, this amplification is suppressed as the value of the differential single-qubit phase ζ becomes comparable to θ . In contrast, MEADD exhibits strong SNR through coherent amplification with no visible dependence on the detuning ζ .

Observe that by going between choosing α or β to be 1 (and the other 0), one engineers the amplitudes α' and β' to take on the values of all the odd-parity cycle matrix elements. By preparing the states $|0+\rangle$ and $|+0\rangle$ and measuring X and Y on each qubit to reconstruct the off-diagonal reduced-density-matrix elements, one obtains the matrix of values $e^{in\phi}C_{\text{odd}}^{2n}$. This protocol is illustrated in Fig. 6(a) and shares some principles with unitary tomography [23]. As pointed out earlier, C_{odd} is a special unitary, with determinant one, and we have

$$\det(e^{in\phi}C_{\text{odd}}^{2n}) = e^{i2n\phi} \det(C_{\text{odd}}^{2n}) = e^{i2n\phi}. \quad (53)$$

This procedure amplifies ϕ coherently $2n$ times, suppressing the variance of our estimates of ϕ by a factor of n^2 . The ability to repeat this measurement for multiple depths n and fit the slope of the accumulated phase makes us robust to SPAM errors. Fluctuations in single-qubit phases are also swept away; identically within the even-parity subspace, and to a large extent within the odd-parity subspace as well. The extent to which they survive will be discussed in more detail when considering the swap angle θ . Decoherence errors are also mitigated, because they tend to affect only the magnitude of the determinant, not its phase. Altogether, this simple protocol yields highly accurate and precise estimates of the controlled phase, as demonstrated experimentally in Fig. 3.

Swap angle and axis. We now show how to estimate θ and χ with the circuits illustrated in Fig. 6(b). The 2-cycle unitary in the odd-parity subspace $\{|01\rangle, |10\rangle\}$

reads

$$C_{\text{odd}}^2 = XU_{\text{odd}}XU_{\text{odd}} \quad (54)$$

$$= Z(-\zeta)e^{-i\theta(\cos\chi X + \sin\chi Y)}Z(-\zeta) \\ \times Z(\zeta)e^{-i\theta(\cos\chi X - \sin\chi Y)}Z(\zeta) \quad (55)$$

$$\approx Z(-\zeta)X(4\theta\cos\chi)Z(\zeta), \quad (56)$$

where the approximation applies for $\theta \ll 1$, which is relevant when implementing controlled-phase gates like CZ (we discuss larger θ to conclude the section). The dynamical decoupling sequence ensures that the only part of the swap that remains is around the X axis in the odd-parity sector. The corresponding cycle Rabi angle $\Omega \approx |\theta\cos\chi|$ can be estimated by observing the population transfer rate between $|01\rangle$ and $|10\rangle$. To determine the swap around the Y axis, we also introduce the complementary cycle unitary

$$\bar{C} = (Y \otimes X)U = \bar{C}_{\text{even}} \oplus \bar{C}_{\text{odd}}, \quad (57)$$

where the second equation is a result of $Y \otimes X = Y \oplus Y$. For $\theta \ll 1$, we have

$$\bar{C}_{\text{odd}}^2 = YU_{\text{odd}}YU_{\text{odd}} \quad (58)$$

$$\approx Z(-\zeta)Y(-4\theta\sin\chi)Z(\zeta), \quad (59)$$

which can be used to estimate $\bar{\Omega} \approx |\theta\sin\chi|$. Again, the variance of the estimates scales as $O(1/n^2)$. Measuring only the population transfer between $|01\rangle$ and $|10\rangle$ leaves $\cos\chi$ and $\sin\chi$ with ambiguous signs. By additionally measuring in the Bell-type basis

$$|\pm\rangle_{\text{odd}} = \frac{1}{\sqrt{2}}(|01\rangle \pm |10\rangle), \quad (60)$$

$$|\pm i\rangle_{\text{odd}} = \frac{1}{\sqrt{2}}(|01\rangle \pm i|10\rangle), \quad (61)$$

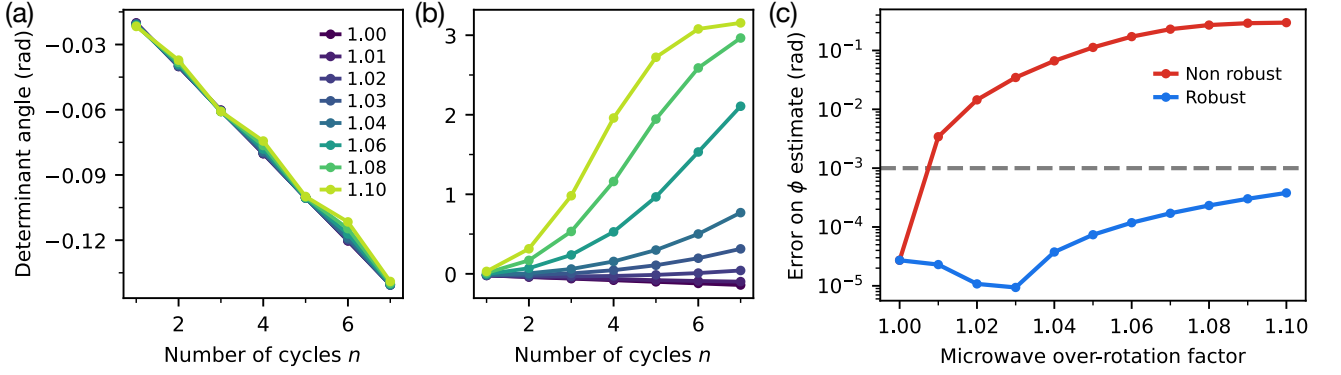


FIG. 8. Numerical simulation results for characterizing the controlled phase of a CZ gate with $\phi = \pi - 0.01$, using MEADD and imperfect DD gates. (a) Using the robust DD sequence described in Sec. V, the angle of the determinant in Eq. (53), which corresponds to $-2n\phi$, builds up linearly even if we introduce microwave over-rotation errors on one of the qubits. The strength of the over-rotation, a factor multiplying the π rotation around X , is varied between 0 and 10% and illustrated with the different colors. (b) Similar simulations now using the XY4 DD sequence, which is *not* robust to implementation errors in this protocol. (c) Resulting ϕ -estimate error obtained from the fitting the slope of the MEADD data for the different over-rotation errors. Dashed line indicates millirad precision.

one resolves these ambiguities and also mitigates the effect of dephasing errors, since one is effectively measuring the complete odd-parity Bloch vector and can distinguish shrinkage from rotation. Since both C and $\bar{C}$ preserve the parity, we can also mitigate single-qubit relaxation errors to first order by postselecting on parity.

For the case of large θ , we can estimate θ and χ using the relations

$$\Omega = \arccos(1 - 2(\sin \theta \cos \chi)^2)/2, \quad (62)$$

$$\bar{\Omega} = \arccos(1 - 2(\sin \theta \sin \chi)^2)/2, \quad (63)$$

where the Rabi angles Ω and $\bar{\Omega}$ can be learned similarly as in the case where $\theta \ll 1$.

In Fig. 7, we compare in simulation MEADD to previously employed techniques for measuring the swap angle: the phase method [5] and unitary tomography [23]. Unitary tomography measures observables whose expectation values correspond to the real and imaginary parts of the matrix elements of a two-qubit, excitation-preserving gate. This is achieved by using the vacuum state as a phase and eigenstate reference. The phase method uses the same initial states and observables as unitary tomography, while repeating the gate to be characterized for various cycle numbers and inserting Z phases at various points to amplify coherent errors and disambiguate different parameter contributions to the composite unitary. As seen in Fig. 7, this coherent amplification gives the phase method higher signal to noise ratio (SNR) than unitary tomography, even when keeping the total number of measurement results constant across both characterization methods. However, one sees that the advantage of coherent amplification is washed out once ζ is on the same order as θ .

In comparison, MEADD exhibits significantly higher SNR and maintains the advantage from coherent amplification even when ζ is significantly larger than θ ,

as the dynamical-decoupling sequence cancels the swap-suppressing effects of ζ .

VI. ROBUSTNESS TO IMPLEMENTATION ERRORS

The dynamical decoupling sequences used by MEADD will unavoidably contain systematic errors, so we must design sequences that are robust to these imperfections. It is well known that the so-called XY4 sequence $(XY)^{2n}$ is robust to any small systematic errors in X and Y [10, 27], whereas the Carr-Purcell-Meiboom-Gill (CPMG) sequence X^{2n} is susceptible to over-rotation errors in X [9]. Similar to XY4, we require our characterization sequence to be robust to any systematic errors in the single-qubit gates. These errors can lead to non-zero matrix elements between the even and odd-parity subspaces, which destroy the direct sum structure in Eq. (44) that we rely on.

In Fig. 8, we demonstrate the importance of using a robust DD sequence to characterize the entangling parameters of a CZ gate with MEADD. We simulate the protocol to estimate the controlled phase detailed in Sec. V, both with the robust CPMG and the non-robust XY4 DD sequences. Using the robust DD sequence in panel (a), we see that the phase builds up linearly even if we introduce microwave over rotation errors on one of the qubits. Even for an over rotation of 10% (0.314 rad), we still get linear phase accumulation to a good approximation that gives us an estimate of ϕ to sub milliradian precision, as shown in panel (c) with the blue curve. On the other hand, using the *non* robust DD sequence gives unusable data as we introduce microwave over rotation errors on one of the qubits, as shown in panel (b). For over rotation of just 2% (0.063 rad), the phase accumulation is so distorted from being linear that the error in the ϕ esti-

mate is larger than the value we are trying to measure (0.01 rad). With this sequence, even 1% over-rotation error introduces errors in the ϕ estimate larger than a milliradian, i.e. it is more than an order of magnitude more sensitive to implementation errors than the robust DD sequence.

In the following, we will first demonstrate how to design robust MEADD protocols for the case of a CZ gate, before moving to general excitation-number conserving two-qubit gates. We consider small arbitrary systematic errors of the X gate as

$$\tilde{X}(\epsilon) = e^{-i(\epsilon_X X + \epsilon_Y Y + \epsilon_Z Z)} X \quad (64)$$

$$\approx (I - i\epsilon \cdot \sigma) X, \quad (65)$$

where $\epsilon = (\epsilon_X, \epsilon_Y, \epsilon_Z)$ and $\sigma = (X, Y, Z)$. Since the errors are assumed to be small, we will consider their effect to first order, which allows us to treat errors on the individual qubits separately in the following.

Robustness for CZ gate. The perfect 2-cycle unitary for $U = \text{CZ}$ reads

$$C^2 = (X \otimes X) \text{CZ} (X \otimes X) \text{CZ} \quad (66)$$

$$= -Z \otimes Z. \quad (67)$$

Noticing that $C^4 = I \otimes I$, we have

$$[(X \otimes X) \text{CZ}]^3 = (Y \otimes Y) \text{CZ} = \text{CZ} (X \otimes X). \quad (68)$$

Considering errors on the right qubit, the noisy 2-cycle unitary is

$$\tilde{C}^2 = (X \otimes \tilde{X}) \text{CZ} (X \otimes \tilde{X}) \text{CZ} \quad (69)$$

$$\approx C^2 - i \left[(I \otimes \epsilon \cdot \sigma) (X \otimes X) \text{CZ} (X \otimes X) \text{CZ} \right. \\ \left. + (X \otimes X) \text{CZ} (I \otimes \epsilon \cdot \sigma) (X \otimes X) \text{CZ} \right], \quad (70)$$

With the identities in Eq. (68), we have

$$\tilde{C}^2 = V C^2, \quad (71)$$

where

$$V = I - i \left[(I \otimes \epsilon \cdot \sigma) + \text{CZ} (I \otimes Y \epsilon \cdot \sigma Y) \text{CZ} \right]. \quad (72)$$

Since X and Y interact with the CZ gate in similar ways, we first consider $\epsilon_X, \epsilon_Y \neq 0$ and $\epsilon_Z = 0$, for which we have the anticommutation relations

$$\{(I \otimes \epsilon \cdot \sigma), Z \otimes Z\} = 0, \quad (73)$$

$$\{\text{CZ} (I \otimes Y \epsilon \cdot \sigma Y) \text{CZ}, Z \otimes Z\} = 0. \quad (74)$$

Both generator terms in the error unitary V in Eq. (72) anticommute with $C^2 = -Z \otimes Z$. Therefore the noisy 4-cycle unitary $\tilde{C}^4 = V C^2 V C^2$ is identical to the noiseless 4-cycle unitary C^4 to first order, and the sequence is robust to small errors in X and Y for a 4-cycle.

For $\epsilon_Z \neq 0$ and $\epsilon_X = \epsilon_Y = 0$, we have

$$\tilde{C}^2 = e^{-i\epsilon_Z (I \otimes Z - I \otimes Z)} C^2 = C^2, \quad (75)$$

and the error ϵ_Z cancels itself in a 2-cycle. The complementary sequence $\tilde{C}$ from Eq. (57) is robust in exactly the same way, the only difference being the Y operators in Eq. (72) maybe replaced by X , which does not affect their anticommutation with $\tilde{C}^2 = -Z \otimes Z$.

One subtlety here is the gauge freedom related to the position assigned to Z errors around microwave gates. It turns out that it can be totally ignored here because one can always commute half of ϵ_Z to the right of the X gate, leaving a pair of $\epsilon_Z/2$ and $-\epsilon_Z/2$ errors conjugating the CZ gate. This gauge has thus no impact at all on the characterization of the controlled phase ϕ or swap angle θ , but is indistinguishable from a change in the swapping phase χ . For the current case of CZ characterization this is acceptable, as for an ideal CZ gate there is no swapping between the qubits ($\theta = 0$), at which point χ has no meaning.

Robustness for crosstalk. We use MEADD to measure the crosstalk (stray coupling) by treating the Hamiltonian Eq. (23) integrated over a given time period as a gate, and executing the dynamical decoupling sequence to measure the small swap angle, exactly as one does for CZ gates. The main difference is that there is no ZZ term as part of the cycle unitary to echo away, so we use XY4 sequences on both qubits, which in this case are robust to microwave over rotation errors because of the lack of a strong entangling gate.

Robustness for arbitrary excitation-preserving gates. Recall from Fig. 2 that the general excitation-preserving gate decomposes into the fundamental entangler $F(\theta, \phi)$ surrounded by Z phases. Defining the Z rotations before and after the fundamental entangler to be K_{pre} and K_{post} , respectively, we get

$$U = K_{\text{post}} F(\theta, \phi) K_{\text{pre}}. \quad (76)$$

Combining the two-qubit gate with the dynamical-decoupling gate, which we notate as D , we represent the repeated cycle unitary as

$$(DU)^n = [DK_{\text{post}} F(\theta, \phi) K_{\text{pre}}]^n \quad (77)$$

$$= K_{\text{pre}}^{-1} [K_{\text{pre}} DK_{\text{post}} F(\theta, \phi)]^n K_{\text{pre}} \quad (78)$$

$$= K_{\text{pre}}^{-1} [D'F]^n K_{\text{pre}}, \quad (79)$$

where we used the gauge freedom to put all the product unitaries after the entangler in the repeated block, and simplified the notation with $D' \equiv K_{\text{pre}} DK_{\text{post}}$ and $F \equiv F(\theta, \phi)$. We can combine both phase fluctuations (which are errors in $K_{\text{pre/post}}$) and microwave imperfections (which are errors in D) into one expression

$$\tilde{D}' \approx (I - i\epsilon E) D', \quad (80)$$

where ϵ is a small parameter and E an error matrix with $\|E\| = O(1)$. To the first order in ϵ , the repeated cycle unitary with error reads

$$[\tilde{D}'F]^n \approx \left(I - i\epsilon \sum_{j=0}^{n-1} [D'F]^j E [D'F]^{-j} \right) [D'F]^n. \quad (81)$$

The condition that the protocol being first-order insensitive to an arbitrary error in $K_{\text{pre/post}}$ or D is

$$\mathcal{T}(E) \equiv \sum_{j=0}^{n-1} [D'F]^j E [D'F]^{-j} = 0, \quad (82)$$

that is, the twirl of E under the group generated by $D'F = K_{\text{pre}} D K_{\text{post}} F(\theta, \phi)$ equals zero. According to our robustness condition, this term must vanish for all E corresponding to single-qubit errors:

$$E \in \text{span}\{I \otimes X, I \otimes Y, I \otimes Z, X \otimes I, Y \otimes I, Z \otimes I\}. \quad (83)$$

The action of conjugating E by $D'F$, called the adjoint action, is an orthogonal map from the 15-dimensional Lie algebra of $\text{SU}(4)$ to itself, which includes the subspace to which E belongs. We notate this map by

$$[\text{Ad}(D'F)](A) = (D'F)A(D'F)^{-1}. \quad (84)$$

As such, we can now analyze the robustness of a dynamical-decoupling sequence by inspecting the eigen-

values of $\text{Ad}(D'F)$. In this new notation,

$$\mathcal{T}(E) = \sum_{j=0}^{n-1} [\text{Ad}(D'F)]^j(E), \quad (85)$$

which is the action of the sum of $[\text{Ad}(D'F)]^j$ on E . Since $\text{Ad}(D'F)$ is an orthogonal map, its eigenvalues are phases. For non trivial phases, this sum will vanish, so we need to ensure that the eigenspaces overlapping with E all have non trivial eigenvalues. We simplify the required analysis by noting that $\text{Ad}(D'F)$ acts trivially on $\text{span}\{X \otimes X, Y \otimes Y, Z \otimes Z\}$. It also respects the $\mathbf{Z}_2$ symmetry of permuting subsystems, so we can further block diagonalize the map into symmetric and antisymmetric subspaces, each 6-dimensional. For representing matrices in these subspaces, we choose for our bases the symmetric/antisymmetric combinations of IX, YZ, IY, ZX, IZ , and XY , in that order.

Using $\text{Ad}_{\pm}$ to denote the adjoint action restricted to these symmetric/antisymmetric combinations, consider the adjoint action of the fundamental entangler:

$$\text{Ad}_+(F(\theta, \phi)) = \begin{pmatrix} \cos(\theta - \phi/2) & -\sin(\theta - \phi/2) & 0 & 0 & 0 & 0 \\ \sin(\theta - \phi/2) & \cos(\theta - \phi/2) & 0 & 0 & 0 & 0 \\ 0 & 0 & \cos(\phi/2 - \theta) & -\sin(\phi/2 - \theta) & 0 & 0 \\ 0 & 0 & \sin(\phi/2 - \theta) & \cos(\phi/2 - \theta) & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \quad (86)$$

$$\text{Ad}_-(F(\theta, \phi)) = \begin{pmatrix} \cos(\theta + \phi/2) & -\sin(\theta + \phi/2) & 0 & 0 & 0 & 0 \\ \sin(\theta + \phi/2) & \cos(\theta + \phi/2) & 0 & 0 & 0 & 0 \\ 0 & 0 & \cos(\theta + \phi/2) & -\sin(\theta + \phi/2) & 0 & 0 \\ 0 & 0 & \sin(\theta + \phi/2) & \cos(\theta + \phi/2) & 0 & 0 \\ 0 & 0 & 0 & 0 & \cos 2\theta & -\sin 2\theta \\ 0 & 0 & 0 & 0 & \sin 2\theta & \cos 2\theta \end{pmatrix}. \quad (87)$$

For the case where the dynamical decoupling layer is $D = X \otimes X$, the adjoint action in both subspaces is diagonal in this basis and respects the 2-dimensional subspaces we just identified:

$$\text{Ad}_{\pm}(D) = \begin{pmatrix} 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & -1 & 0 & 0 & 0 \\ 0 & 0 & 0 & -1 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 & 0 & -1 \end{pmatrix}. \quad (88)$$

The combination of CZ with the dynamical decoupling

gives symmetric and antisymmetric actions

$$\text{Ad}_+(DCZ) = \begin{pmatrix} 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & -1 & 0 & 0 & 0 & 0 \\ -1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}, \quad (89)$$

$$\text{Ad}_-(DCZ) = \begin{pmatrix} 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & -1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ -1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 & 0 & -1 \end{pmatrix}. \quad (90)$$

The eigenvalues are non trivial with values $-1, i$, and $-i$, except for the last diagonal entry in Ad_+ , which corre-

sponds to the symmetric combination of XY. Since this is a non-local, non excitation preserving error that is not expected to be present in our system, we do not concern ourselves with echoing this away. Taking sums of powers of the other eigenvalues, $\sum_j \lambda^j$, we find that all vanish when including four terms in the sum, which agrees with our earlier analysis, where the 4-cycle unitary was identical with the noiseless 4-cycle unitary when considering all single-qubit errors.

We can carry this same analysis out for other fSim gates. Consider $\sqrt{\text{iSWAP}} = F(\pi/4, 0)$, for instance:

$$\begin{aligned} \text{Ad}_+(D\sqrt{\text{iSWAP}}) \\ = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & -1 & 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & -1 & -1 & 0 & 0 \\ 0 & 0 & 1 & -1 & 0 & 0 \\ 0 & 0 & 0 & 0 & -\sqrt{2} & 0 \\ 0 & 0 & 0 & 0 & 0 & -\sqrt{2} \end{pmatrix}, \end{aligned} \quad (91)$$

$$\begin{aligned} \text{Ad}_-(D\sqrt{\text{iSWAP}}) \\ = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & -1 & 0 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & -1 & 1 & 0 & 0 \\ 0 & 0 & -1 & -1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & \sqrt{2} \\ 0 & 0 & 0 & 0 & -\sqrt{2} & 0 \end{pmatrix}. \end{aligned} \quad (92)$$

None of these eigenvalues are trivial, so everything is echoed away. Because the phases are smaller, however, it takes 8 cycles instead of 4 before complete cancellation occurs, which puts stricter limits on how noisy the microwaves can be before the errors start polluting the characterization.

Finally, consider the iSWAP gate itself, $F(\pi/2, 0)$. Looking at its adjoint action, we notice a problem:

$$\text{Ad}_+(\text{iSWAP}) = \begin{pmatrix} 0 & -1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & -1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix} \quad (93)$$

$$\text{Ad}_-(\text{iSWAP}) = \begin{pmatrix} 0 & -1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & -1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 & 0 & -1 \end{pmatrix}. \quad (94)$$

Since the symmetric (IZ, XY) subspace is degenerate with trivial eigenvalue +1, we need to use the dynamical-decoupling gates. However, the antisymmetric (IZ, XY) subspace is also degenerate, with eigenvalue -1, using the dynamical-decoupling gate will flip that eigenvalue to +1, which will allow antisymmetric Z errors to be amplified.

One strategy for overcoming this problem is to only use the iSWAP gate in every other cycle. In the cycles where we do not use the iSWAP gate, we will instead idle the qubits for the iSWAP gate time. This ensures the Z phase errors are identical between cycles, and thus coherently cancel with one another. Since adjacent ideal $X \otimes X$ gates cancel, the sum of adjoint actions we need to consider for this sequence is

$$\begin{aligned} \sum_j \text{Ad}_\pm(\text{iSWAP}^j) + \text{Ad}_\pm(D\text{iSWAP}^j) \\ = (1 + \text{Ad}_\pm(D)) \sum_j \text{Ad}_\pm(\text{iSWAP})^j. \end{aligned} \quad (95)$$

The prefactor $1 + \text{Ad}_\pm(D)$ annihilates all but the (IX, YZ) subspace (symmetric and antisymmetric), and we see that iSWAP effectively echos away those errors on its own, so this sequence echos away low-frequency Z noise and is robust to imperfections in the microwave gates.

VII. CONCLUSION

We have introduced a unitary characterization framework called MEADD that significantly surpasses existing approaches for frequency-tunable transmons, both in terms of accuracy and precision, while being experimentally efficient. It leverages the noise structure and control precision hierarchy of the quantum device to yield gate characterizations that are largely insensitive to qubit frequency noise and robust to imperfections in the extra dynamical decoupling gates required for its implementation. With MEADD, we have experimentally demonstrated precise measurements of both single- and two-qubit gate parameters and crosstalk. Remarkably, MEADD can readily be used to estimate the entangling parameters of CZ gates with 5x and 10x higher precision than previous methods for the controlled-phase and swap angle, respectively. This characterization information readily identifies systematic errors present in the implemented quantum operations which can then be corrected for by adjusting the control parameters. We expect that insights making MEADD useful in characterizing superconducting qubits can be applied to other qubit platforms, including leveraging the most precise control parameters for estimating less precise ones, isolating stable parameters from less stable and noisy ones, and targeting only a few parameters at a time.

VIII. ACKNOWLEDGEMENTS

We are grateful to the Google Quantum AI team for building, operating, and maintaining software and hardware infrastructure used in this work. The authors would like to thank Abraham Asfaw, Sergio Boixo, Yu Chen, Ben Chiaro, Michel Devoret, Vinicius S. Ferreira, Evan Jeffrey, Amir Karamlou, Julian Kelly, Paul V. Klimov,

Alexander Korotkov, Kenny Lee, Will Livingston, Xiao Mi, Charles Neill, Murphy Yuezhen Niu, Will Oliver, Andre Petukhov, and Vadim N. Smelyanskiy for discus-

sions and feedback on the draft, and Andreas Bengtsson, Catherine Erickson, and Sabrina Hong for software support.

-
- [1] Robin Blume-Kohout, John King Gamble, Erik Nielsen, Jonathan Mizrahi, Jonathan D. Sterk, and Peter Maunz, “Robust, self-consistent, closed-form tomography of quantum logic gates on a trapped ion qubit,” (2013), [arXiv:1310.4492 \[quant-ph\]](#).
 - [2] Robin Blume-Kohout, John King Gamble, Erik Nielsen, Kenneth Rudinger, Jonathan Mizrahi, Kevin Fortier, and Peter Maunz, “Demonstration of qubit operations below a rigorous fault tolerance threshold with gate set tomography,” *Nature Communications* **8**, 1–13 (2017), number: 1 Publisher: Nature Publishing Group.
 - [3] Shelby Kimmel, Guang Hao Low, and Theodore J. Yoder, “Robust calibration of a universal single-qubit gate set via robust phase estimation,” *Physical Review A* **92**, 062315 (2015), publisher: American Physical Society.
 - [4] Frank Arute, Kunal Arya, Ryan Babbush, Dave Bacon, Joseph C. Bardin, Rami Barends, Andreas Bengtsson, Sergio Boixo, Michael Broughton, Bob B. Buckley, David A. Buell, Brian Burkett, Nicholas Bushnell, Yu Chen, Zijun Chen, Yu-An Chen, Ben Chiaro, Roberto Collins, Stephen J. Cotton, William Courtney, Sean Demura, Alan Derk, Andrew Dunsworth, Daniel Eppens, Thomas Ekl, Catherine Erickson, Edward Farhi, Austin Fowler, Brooks Foxen, Craig Gidney, Marissa Giustina, Rob Graff, Jonathan A. Gross, Steve Habegger, Matthew P. Harrigan, Alan Ho, Sabrina Hong, Trent Huang, William Huggins, Lev B. Ioffe, Sergei V. Isakov, Evan Jeffrey, Zhang Jiang, Cody Jones, Dvir Kafri, Kostyantyn Kechedzhi, Julian Kelly, Seon Kim, Paul V. Klimov, Alexander N. Korotkov, Fedor Kostritsa, David Landhuis, Pavel Laptev, Mike Lindmark, Erik Lucero, Michael Marthaler, Orion Martin, John M. Martinis, Anika Maruszyk, Sam McArdle, Jarrod R. McClean, Trevor McCourt, Matt McEwen, Anthony Megrant, Carlos Mejuto-Zaera, Xiao Mi, Masoud Mohseni, Wojciech Mruczkiewicz, Josh Mutus, Ofer Naaman, Matthew Neeley, Charles Neill, Hartmut Neven, Michael Newman, Murphy Yuezhen Niu, Thomas E. O’Brien, Eric Ostby, Bálint Pató, Andre Petukhov, Harald Putterman, Chris Quintana, Jan-Michael Reiner, Pedram Roushan, Nicholas C. Rubin, Daniel Sank, Kevin J. Satzinger, Vadim Smelyanskiy, Doug Strain, Kevin J. Sung, Peter Schmitteckert, Marco Szalay, Norm M. Tubman, Amit Vainsencher, Theodore White, Nicolas Vogt, Z. Jamie Yao, Ping Yeh, Adam Zalcman, and Sebastian Zanker, “Observation of separated dynamics of charge and spin in the Fermi-Hubbard model,” [arXiv:2010.07965 \[quant-ph\]](#) (2020).
 - [5] C. Neill, T. McCourt, X. Mi, Z. Jiang, M. Y. Niu, W. Mruczkiewicz, I. Aleiner, F. Arute, K. Arya, J. Atalaya, R. Babbush, J. C. Bardin, R. Barends, A. Bengtsson, A. Bourassa, M. Broughton, B. B. Buckley, D. A. Buell, B. Burkett, N. Bushnell, J. Campero, Z. Chen, B. Chiaro, R. Collins, W. Courtney, S. Demura, A. R. Derk, A. Dunsworth, D. Eppens, C. Erickson, E. Farhi, A. G. Fowler, B. Foxen, C. Gidney, M. Giustina, J. A. Gross, M. P. Harrigan, S. D. Harrington, J. Hilton, A. Ho, S. Hong, T. Huang, W. J. Huggins, S. V. Isakov, M. Jacob-Mitos, E. Jeffrey, C. Jones, D. Kafri, K. Kechedzhi, J. Kelly, S. Kim, P. V. Klimov, A. N. Korotkov, F. Kostritsa, D. Landhuis, P. Laptev, E. Lucero, O. Martin, J. R. McClean, M. McEwen, A. Megrant, K. C. Miao, M. Mohseni, J. Mutus, O. Naaman, M. Neeley, M. Newman, T. E. O’Brien, A. Opremcak, E. Ostby, B. Pató, A. Petukhov, C. Quintana, N. Redd, N. C. Rubin, D. Sank, K. J. Satzinger, V. Shvarts, D. Strain, M. Szalay, M. D. Trevithick, B. Villalonga, T. C. White, Z. Yao, P. Yeh, A. Zalcman, H. Neven, S. Boixo, L. B. Ioffe, P. Roushan, Y. Chen, and V. Smelyanskiy, “Accurately computing the electronic properties of a quantum ring,” *Nature* **594**, 508–512 (2021).
 - [6] I. Chiorescu, Y. Nakamura, C. J. P. M. Harmans, and J. E. Mooij, “Coherent Quantum Dynamics of a Superconducting Flux Qubit,” *Science* **299**, 1869–1871 (2003).
 - [7] C. M. Quintana, Yu Chen, D. Sank, A. G. Petukhov, T. C. White, Dvir Kafri, B. Chiaro, A. Megrant, R. Barends, B. Campbell, Z. Chen, A. Dunsworth, A. G. Fowler, R. Graff, E. Jeffrey, J. Kelly, E. Lucero, J. Y. Mutus, M. Neeley, C. Neill, P. J. J. O’Malley, P. Roushan, A. Shabani, V. N. Smelyanskiy, A. Vainsencher, J. Wenner, H. Neven, and John M. Martinis, “Observation of classical-quantum crossover of $1/f$ flux noise and its paramagnetic temperature dependence,” *Phys. Rev. Lett.* **118**, 057702 (2017).
 - [8] H. Y. Carr and E. M. Purcell, “Effects of diffusion on free precession in nuclear magnetic resonance experiments,” *Phys. Rev.* **94**, 630–638 (1954).
 - [9] S. Meiboom and D. Gill, “Modified Spin-Echo Method for Measuring Nuclear Relaxation Times,” *Review of Scientific Instruments* **29**, 688–691 (1958).
 - [10] Lorenza Viola, Emanuel Knill, and Seth Lloyd, “Dynamical decoupling of open quantum systems,” *Phys. Rev. Lett.* **82**, 2417–2421 (1999).
 - [11] K. Khodjasteh and D. A. Lidar, “Fault-tolerant quantum dynamical decoupling,” *Phys. Rev. Lett.* **95**, 180501 (2005).
 - [12] Götz S. Uhrig, “Keeping a quantum bit alive by optimized π -pulse sequences,” *Phys. Rev. Lett.* **98**, 100504 (2007).
 - [13] Frank Arute, Kunal Arya, Ryan Babbush, Dave Bacon, Joseph C. Bardin, Rami Barends, Rupak Biswas, Sergio Boixo, Fernando G. S. L. Brandao, David A. Buell, Brian Burkett, Yu Chen, Zijun Chen, Ben Chiaro, Roberto Collins, William Courtney, Andrew Dunsworth, Edward Farhi, Brooks Foxen, Austin Fowler, Craig Gidney, Marissa Giustina, Rob Graff, Keith Guerin, Steve Habegger, Matthew P. Harrigan, Michael J. Hartmann, Alan Ho, Markus Hoffmann, Trent Huang, Travis S. Humble, Sergei V. Isakov, Evan Jeffrey, Zhang Jiang, Dvir Kafri, Kostyantyn Kechedzhi, Julian Kelly, Paul V. Klimov, Sergey Knysh, Alexander Korotkov, Fedor Kostritsa, David Landhuis, Mike Lindmark, Erik Lucero, Dmitry

- Lyakh, Salvatore Mandrà, Jarrod R. McClean, Matthew McEwen, Anthony Megrant, Xiao Mi, Kristel Michielsen, Masoud Mohseni, Josh Mutus, Ofer Naaman, Matthew Neeley, Charles Neill, Murphy Yuezhen Niu, Eric Ostby, Andre Petukhov, John C. Platt, Chris Quintana, Eleanor G. Rieffel, Pedram Roushan, Nicholas C. Rubin, Daniel Sank, Kevin J. Satzinger, Vadim Smelyanskiy, Kevin J. Sung, Matthew D. Trevithick, Amit Vainsencher, Benjamin Villalonga, Theodore White, Z. Jamie Yao, Ping Yeh, Adam Zalcman, Hartmut Neven, and John M. Martinis, “Quantum supremacy using a programmable superconducting processor,” *Nature* **574**, 505–510 (2019).
- [14] Sheng-Tao Wang, Dong-Ling Deng, and L.-M. Duan, “Hamiltonian tomography for quantum many-body systems with arbitrary couplings,” *New Journal of Physics* **17**, 093017 (2015), publisher: IOP Publishing.
- [15] Yulong Dong, Jonathan Gross, and Murphy Yuezhen Niu, “Beyond Heisenberg limit quantum metrology through quantum signal processing,” (2022), arXiv:2209.11207 [quant-ph].
- [16] Uri Vool and Michel Devoret, “Introduction to quantum electromagnetic circuits,” *International Journal of Circuit Theory and Applications* **45**, 897–934 (2017).
- [17] Philip Krantz, Morten Kjaergaard, Fei Yan, Terry P. Orlando, Simon Gustavsson, and William D. Oliver, “A quantum engineer’s guide to superconducting qubits,” *Applied Physics Reviews* **6**, 021318 (2019).
- [18] Alexandre Blais, Arne L. Grimsmo, S. M. Girvin, and Andreas Wallraff, “Circuit Quantum Electrodynamics,” *Reviews of Modern Physics* **93**, 025005 (2021).
- [19] F. Motzoi, J. M. Gambetta, P. Rebentrost, and F. K. Wilhelm, “Simple pulses for elimination of leakage in weakly nonlinear qubits,” *Physical Review Letters* **103**, 110501 (2009), arXiv:0901.0534 [cond-mat, physics:quant-ph].
- [20] M. A. Rol, F. Battistel, F. K. Malinowski, C. C. Bultink, B. M. Tarasinski, R. Vollmer, N. Haider, N. Muthusubramanian, A. Bruno, B. M. Terhal, and L. DiCarlo, “Fast, high-fidelity conditional-phase gate exploiting leakage interference in weakly anharmonic superconducting qubits,” *Phys. Rev. Lett.* **123**, 120502 (2019).
- [21] V. Negîrneac, H. Ali, N. Muthusubramanian, F. Battistel, R. Sagastizabal, M. S. Moreira, J. F. Marques, W. J. Vlothuizen, M. Beekman, C. Zachariadis, N. Haider, A. Bruno, and L. DiCarlo, “High-fidelity controlled-z gate with maximal intermediate leakage operating at the speed limit in a superconducting quantum processor,” *Phys. Rev. Lett.* **126**, 220502 (2021).
- [22] B.R. Johnson, “Controlling Photons in Superconducting Electrical Circuits,” Ph.D. Dissertation, Yale University, 2011.
- [23] B. Foxen, C. Neill, A. Dunsworth, P. Roushan, B. Chiaro, A. Megrant, J. Kelly, Zijun Chen, K. Satzinger, R. Barends, F. Arute, K. Arya, R. Babbush, D. Bacon, J. C. Bardin, S. Boixo, D. Buell, B. Burkett, Yu Chen, R. Collins, E. Farhi, A. Fowler, C. Gidney, M. Giustina, R. Graff, M. Harrigan, T. Huang, S. V. Isakov, E. Jeffrey, Z. Jiang, D. Kafri, K. Kechedzhi, P. Klimov, A. Korotkov, F. Kostritsa, D. Landhuis, E. Lucero, J. McClean, M. McEwen, X. Mi, M. Mohseni, J. Y. Mutus, O. Naaman, M. Neeley, M. Niu, A. Petukhov, C. Quintana, N. Rubin, D. Sank, V. Smelyanskiy, A. Vainsencher, T. C. White, Z. Yao, P. Yeh, A. Zalcman, H. Neven, and J. M. Martinis (Google AI Quantum), “Demonstrating a continuous set of two-qubit gates for near-term quantum algorithms,” *Phys. Rev. Lett.* **125**, 120504 (2020).
- [24] Byron Drury and Peter J. Love, “Constructive quantum Shannon decomposition from Cartan involutions,” *Journal of Physics A: Mathematical and Theoretical* **41**, 395305 (2008), arXiv:0806.4015 [quant-ph].
- [25] Dripto M. Debroy, Élie Genois, Jonathan A. Gross, Wojciech Mruczkiewicz, Kenny Lee, Sabrina Hong, Zijun Chen, Vadim Smelyanskiy, and Zhang Jiang, “Context-aware fidelity estimation,” *Phys. Rev. Res.* **5**, 043202 (2023).
- [26] “Suppressing quantum errors by scaling a surface code logical qubit,” *Nature* **614**, 676–681 (2023).
- [27] A. A. Maudsley, “Modified Carr-Purcell-Meiboom-Gill sequence for NMR fourier imaging applications,” *Journal of Magnetic Resonance* (1969) **69**, 488–491 (1986).

Appendix A: Derivative Removal by Adiabatic Gate (DRAG)

In this appendix, we present a simple derivation of the DRAG pulse used to mitigate gate leakage [19]. We model a superconducting qubit as a driven nonlinear oscillator similar to Eq. (3) in the main text ($\hbar = 1$),

$$H_I(t) = \delta\omega_{01}(t) a^\dagger a - \frac{\eta}{2} a^{\dagger 2} a^2 + [\Omega(t) a^\dagger + \text{h.c.}], \quad (\text{A1})$$

where $a^\dagger$ (a) is the bosonic creation (annihilation) operators and η the so-called anharmonicity. Setting the detuning $\delta\omega_{01}(t) = 0$ and truncating the Hamiltonian to the second level $|2\rangle$, we have

$$H_I(t) = \begin{pmatrix} 0 & \Omega(t)^* & 0 \\ \Omega(t) & 0 & \sqrt{2}\Omega(t)^* \\ 0 & \sqrt{2}\Omega(t) & -\eta \end{pmatrix}. \quad (\text{A2})$$

The time-dependent Schrödinger equation under $H_I(t)$ reads

$$i \frac{\partial}{\partial t} |\psi_I(t)\rangle = H_I(t) |\psi_I(t)\rangle. \quad (\text{A3})$$

Going to a second interaction picture $|\psi_{II}(t)\rangle = e^{-i\eta t |2\rangle\langle 2|} |\psi_I(t)\rangle$, we have

$$i \frac{\partial}{\partial t} |\psi_{II}(t)\rangle = H_{II}(t) |\psi_{II}(t)\rangle, \quad (\text{A4})$$

where the off-diagonal Hamiltonian H_{II} reads

$$H_{II}(t) = \begin{pmatrix} 0 & \Omega(t)^* & 0 \\ \Omega(t) & 0 & \sqrt{2}\Omega(t)^* e^{i\eta t} \\ 0 & \sqrt{2}\Omega(t) e^{-i\eta t} & 0 \end{pmatrix}. \quad (\text{A5})$$

Using first-order time-dependent perturbation theory, we obtain the leakage amplitude at the end of the microwave

pulse for $-T/2 \leq t \leq T/2$,

$$\langle 2 | \psi_H(T/2) \rangle \approx -i \int_{-\frac{T}{2}}^{\frac{T}{2}} \langle 2 | H_H(t) | 1 \rangle \langle 1 | \psi_H(t) \rangle dt \quad (\text{A6})$$

$$= \int_{-\frac{T}{2}}^{\frac{T}{2}} L(t) e^{-i\eta t} dt \quad (\text{A7})$$

$$\equiv \tilde{L}(\eta), \quad (\text{A8})$$

where $L(t) \equiv -i\sqrt{2}\Omega(t)\langle 1 | \psi_H(t) \rangle$ and $\tilde{L}(\eta)$ its Fourier transformation. For any differentiable $L(t)$, we have

$$\tilde{L}(\eta) = \frac{i}{\eta} \int_{-\frac{T}{2}}^{\frac{T}{2}} L(t) de^{-i\eta t} \quad (\text{A9})$$

$$= \frac{i}{\eta} L(t) e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}} + \frac{i}{\eta} \int_{-\frac{T}{2}}^{\frac{T}{2}} L^{(1)}(t) e^{-i\eta t} dt, \quad (\text{A10})$$

where $L^{(n)}(t) = d^n L(t)/dt^n$. Using integration by parts again, we have

$$\begin{aligned} \tilde{L}(\eta) \approx & \frac{i}{\eta} L(t) e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}} - \frac{1}{\eta^2} L^{(1)}(t) e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}} \\ & - \frac{i}{\eta^3} L^{(2)}(t) e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}}, \end{aligned} \quad (\text{A11})$$

where higher order terms are truncated when $|L^{(n+1)}(t)/L^{(n)}(t)| \ll \eta$. We choose a control pulse $\Omega(t)$ whose value and first derivative vanish at both boundaries,

$$\Omega(-T/2) = \Omega(T/2) = 0, \quad (\text{A12})$$

$$\Omega^{(1)}(-T/2) = \Omega^{(1)}(T/2) = 0. \quad (\text{A13})$$

Putting Eq. (A12) into the Schrödinger equation (A4), we have

$$\langle 1 | \psi_H^{(1)}(-T/2) \rangle = \langle 1 | \psi_H^{(1)}(T/2) \rangle = 0. \quad (\text{A14})$$

With identities (A12)-(A14), the first two terms in Eq. (A11) vanish, and we can express the leakage amplitude as

$$\tilde{L}(\eta) \approx -\frac{i}{\eta^3} L^{(2)}(t) e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}} \quad (\text{A15})$$

$$= -\frac{\sqrt{2}}{\eta^3} \Omega^{(2)}(t) \langle 1 | \psi_H(t) \rangle e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}}. \quad (\text{A16})$$

To zero out the leakage, we add a derivative term to the original pulse $\Omega(t)$,

$$\Omega_{\text{DRAG}}(t) = \Omega(t) - \frac{i}{\eta} \Omega^{(1)}(t), \quad (\text{A17})$$

and correspondingly, we have

$$L_{\text{DRAG}}(t) = L(t) - \frac{\sqrt{2}}{\eta} \Omega^{(1)}(t) \langle 1 | \psi_H(t) \rangle. \quad (\text{A18})$$

Replacing $L(t)$ with $L_{\text{DRAG}}(t)$ in Eq. (A11) leads to the leakage amplitude for the DRAG pulse. The leading contribution from the added derivative term in $L_{\text{DRAG}}(t)$ comes from the term with coefficient $1/\eta^2$ in Eq. (A11)

$$\begin{aligned} & -\frac{1}{\eta^2} \frac{d}{dt} \left(-\frac{\sqrt{2}}{\eta} \Omega^{(1)}(t) \langle 1 | \psi_H(t) \rangle \right) e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}} \\ & = \frac{\sqrt{2}}{\eta^3} \Omega^{(2)}(t) \langle 1 | \psi_H(t) \rangle e^{-i\eta t} \Big|_{-\frac{T}{2}}^{\frac{T}{2}}, \end{aligned} \quad (\text{A19})$$

which exactly cancels Eq. (A16) and eliminates gate leakages.

Appendix B: Floquet characterization for two-qubit gates

For complete characterization of the two-qubit gate, we need the single-qubit phases ζ and γ in addition to the entangling parameters θ , χ , and ϕ . To accomplish this, we use a method very similar to Floquet characterization [4]. We are able to simplify the protocol by removing the variable Z rotations since we can use our precise estimate of θ to infer ζ from the energy splitting in the odd-parity sector.

We repeat the unitary for many different depths n , using preparation and measurement analogous to what is used to arrive at Eqs. (51) and (52) and shown in Fig. 6, only this time we are directly measuring the matrix elements of the gate unitary U_{odd} instead of C_{odd} , which is interleaved with DD gates. Computing the complex angle of the determinant at each depth n allows us extract γ from the slope exactly as we extracted ϕ from the slope in Eq. (53).

To extract ζ , we use another technique for mitigating the effects of decoherence. The singular value decomposition of the experimentally measured matrix M (which in the ideal case is simply U_{odd}) give us

$$M = U_M \Sigma_M V_M^\dagger. \quad (\text{B1})$$

Just as we took the complex angle of the determinant to mitigate decoherence effects, we will take the unitary part of M to mitigate decoherence effects:

$$M' = U_M V_M^\dagger \quad (\text{B2})$$

In the case of interest where swapping is small, the eigenstates of M' can be unambiguously associated with $|01\rangle$ and $|10\rangle$, allowing us to consistently define the sign as well as the magnitude of the rotation angle of this unitary. Fitting the rotation angle as a function of depth n allows us to extract the rotation angle per gate Ω . Finally, we extract ζ using the relation $\cos \Omega = \cos \zeta \cos \theta$ and $\text{sgn } \zeta = \text{sgn } \Omega$:

$$\zeta = \text{sgn } \Omega \arccos(\cos \Omega / \cos \theta). \quad (\text{B3})$$

There is a danger here that noise causes $\cos \Omega$ to exceed $\cos \theta$, in which case we set $\zeta = 0$.

This method leverages the amplification and SPAM robustness of repeating the unitary to measure γ and Ω , while making use of our precise measurement of θ to ultimately extract ζ (something that had to be accomplished with additional experiments using physical Z rotations in Floquet [4] and Phase-method [5] characterizations).

Appendix C: Simulation parameters

In Fig. 7, for each of the 100 executions at a particular parameter point, unitary tomography and phase method collected 160,000 measurement results. Since the phase method uses four different repetition numbers (1, 2, 4,

and 8), this meant only collecting 2,500 samples per circuit as opposed to 10,000 samples per circuit for unitary tomography. Phase method uses two different Z rotations for each repetition number, without doing measurement flipping, while unitary tomography does not add any extra Z phases, but does do measurement flipping, so at each depth the number of circuits is the same. We did not simulate realistic asymmetric readout noise, so the measurement flipping is essentially just giving twice as many samples at that point. The MEADD circuits used in Fig. 7 only measures two of the five parameters (θ and χ), so the number of measurement results were lowered to 64,000.

5.5 Perspectives

En plus de l'application des protocoles CAFE et MEADD à d'autres opérations (par exemple des portes qui ne font pas partie du groupe de Clifford ou des opérations à plus de deux qubits) et à d'autres plateformes de processeurs quantiques (par exemple des ions piégés ou des atomes neutres), il existe plusieurs adaptations et extensions dont le développement pourrait s'avérer particulièrement utile.

Tout d'abord pour CAFE, le modèle utilisé pour approcher les données et extraire le bilan des contributions à l'erreur moyenne totale pourrait être grandement raffiné. Plus spécifiquement, nous utilisons un simple canal de dépolarisation pour décrire les processus incohérents, alors que nous savons que des modèles microscopiques du bruit incluant de la relaxation et du déphasage décrivent mieux les dynamiques des portes logiques. Il serait alors intéressant d'utiliser différents modèles afin d'approcher les données produites par CAFE et de caractériser quelles dynamiques sont significatives et d'en apprendre plus sur le comportement réel de l'opération réalisée, une méthode qui s'apparente encore une fois à la sélection de modèles [161]. D'ailleurs, plutôt que d'ajuster la moyenne des circuits de CAFE, il serait possible d'utiliser l'information plus granulaire de chacun des circuits formant le 2 -design (4^n pour n qubits) afin d'extraire davantage d'information sur les processus cohérents et sur la non-uniformité du bruit stochastique.

Tirant profit du fait que CAFE offre une excellente visibilité sur les erreurs cohérentes associées à une opération, il pourrait être grandement bénéfique d'utiliser cette métrique dans l'optimisation d'opérations en boucle fermée sur le dispositif. En effet, l'utilisation de CAFE permettrait à la fonction de coût de quantifier la performance des opérations testées de manière plus détaillée, ce qui pourrait améliorer la convergence et réduire les ressources expérimentales nécessaires à l'optimisation d'opération utilisant directement la fidélité fournie par une approche d'étalonnage par portes aléatoires [116, 117, 118, 151].

Grâce à son estimation précise et exacte des paramètres unitaires de portes logiques à un et deux qubits, MEADD peut également être utilisée directement dans la calibration de ces opérations, par exemple en construisant la fonction de coût à partir de l'information fournie par MEADD dans une boucle d'optimisation basée sur l'apprentissage par renforcement. Ainsi, au lieu ou en plus d'optimiser simplement la fidélité de porte moyenne, la fonction de coût d'optimisation pourrait prendre la forme quadratique suivante

$$\mathcal{C}_{\text{MEADD}} = w_{\theta}(\tilde{\theta} - \bar{\theta})^2 + w_{\phi}(\tilde{\phi} - \bar{\phi})^2, \quad (5.16)$$

où $(\tilde{\theta}, \tilde{\phi})$ sont les estimations expérimentales fournies par MEADD des paramètres de la

porte fSim réalisée expérimentalement, $(\bar{\theta}, \bar{\phi})$ sont les paramètres cibles (par exemple $\bar{\theta} = 0$ et $\bar{\phi} = \pi$ pour un CZ) et w_j sont des préfacteurs permettant d'ajuster si l'un des paramètres unitaires est à prioriser dans l'optimisation.

Enfin, MEADD pourrait être généralisé à la caractérisation d'opérations unitaires impliquant davantage de qubits afin de caractériser et éventuellement calibrer un groupe élargi d'opérations. Par exemple, la mesure des stabilisateurs dans le code de surface requiert une transformation unitaire impliquant 5 qubits et 4 portes CZ (en plus de quelques portes à un qubit que l'on pourrait potentiellement négliger en raison de leur fidélité largement supérieure). Plusieurs imperfections peuvent faire en sorte que l'unitaire associée à la réalisation de ces stabilisateurs ne correspond pas directement à la combinaison des 4 unitaires à deux qubits. Étant donné que les erreurs cohérentes associées à la réalisation imparfaite de la mesure des stabilisateurs peuvent être particulièrement dommageables pour le code de surface [184], il serait très intéressant d'utiliser MEADD pour caractériser avec précision ces erreurs.

Chapitre 6

Optimisation de la mesure en circuit QED

Au chapitre 4, nous avons développé une approche combinant l'apprentissage automatique à des techniques existantes de contrôle optimal quantique afin de réaliser des portes logiques de haute fidélité à l'aide de qubits supraconducteurs. Ce chapitre porte également sur le contrôle optimal quantique, mais on s'intéresse ici à la résolution de problèmes de contrôle où les approches existantes sont inadéquates. Plus spécifiquement, on cherche à contrôler des systèmes quantiques complexes où la dissipation joue un rôle clé qu'on ne peut négliger. Alors que les solutions existantes nécessitent des ressources numériques prohibitives pour résoudre ces problèmes, nous proposons une approche alternative permettant d'obtenir les contrôles optimaux pour tout système où la dynamique quantique peut être décrite par une équation maîtresse sous forme de Lindblad. Afin de démontrer l'utilité et le potentiel de cette approche, nous l'appliquons à la résolution de deux opérations de grande importance ayant typiquement des fidélités moindres que les portes logiques unitaires : la mesure et la réinitialisation de transmons.

Dans les sections qui suivent, j'introduis l'importance ainsi que les problématiques associées au contrôle optimal quantique de systèmes ouverts, avant de présenter à la section 6.2 une vue d'ensemble sur la méthodologie que nous avons développée ainsi que sur les résultats que nous avons obtenus en l'appliquant au contrôle d'un transmon. Après la présentation de l'article scientifique portant sur cette recherche à la section 6.3, je discute à la section 6.4 de quelques directions de recherche prometteuses en plus d'introduire brièvement la librairie numérique que nous avons développée en lien avec ce projet : `DYNAMICS` [79].

6.1 Mise en contexte

La généralisation du contrôle optimal quantique à partir d'opérations purement unitaires, telles que le transfert d'un état quantique à un autre et la réalisation de portes logiques, vers l'optimisation d'opérations sur des systèmes quantiques ouverts est intéressante pour plusieurs raisons. Tout d'abord, plusieurs opérations nécessaires au traitement de l'information quantique sont décrites par des processus dissipatifs, notamment la mesure et l'initialisation d'états quantiques. Ainsi, l'optimisation des contrôles réalisant de telles opérations requière une description des dynamiques quantiques qui est plus générale qu'une transformation unitaire, en l'occurrence l'utilisation d'une transformation complètement positive et préservant la trace de la matrice densité [48]. Comme décrit au chapitre 2, la représentation la plus répandue en circuits supraconducteurs d'une telle transformation est une équation maîtresse sous forme de Lindblad. Cette représentation considère que le système d'intérêt est couplé faiblement à un environnement dont les dynamiques s'atténuent beaucoup plus rapidement que l'échelle de temps des dynamiques du système. Pour les qubits supraconducteurs ainsi que plusieurs autres plateformes, un aspect rendant le contrôle optimal en système ouvert particulièrement prometteur est le fait que certaines dynamiques dissipatives sont contrôlables, ouvrant ainsi la porte aux nombreuses applications connues en ingénierie de la dissipation quantique [185]. Par exemple, les transmons sont typiquement couplés à une cavité micro-onde, laquelle est dissipative en raison de son couplage important avec une ligne de transmission. Il est alors possible d'appliquer des impulsions externes sur le qubit ou sur la cavité afin de modifier la cohérence du transmon, par exemple en tirant profit de l'interaction dispersive qui est également utilisée pour sa mesure.

L'inclusion explicite des processus dissipatifs dans l'optimisation peut également mener à de meilleures opérations unitaires en exploitant des sous-espaces qui sont moins affectés par la décohérence associée à l'environnement [186]. L'exemple par excellence illustrant le possible avantage d'une telle approche est la présence de sous-espaces sans décohérence dans un système où il existe une symétrie entre le système et son environnement, par exemple lorsque plusieurs qubits sont couplés de façon indiscernable au même environnement [187]. Je note enfin que dans plusieurs situations où l'environnement possède une densité spectrale de bruit non uniforme ou des temps de corrélation non négligeables avec une certaine mémoire de durée finie, il est possible d'utiliser une approche de contrôle optimal en systèmes quantiques ouverts afin d'exploiter cette structure et de réaliser de meilleures opérations quantiques. L'utilisation de contrôles de découplage dynamique (DD) pour allonger la cohérence de l'information quantique en est un parfait exemple [174].

Malheureusement, la résolution numérique du contrôle optimal d'un système quantique

ouvert est significativement plus ardue que celle d'un système fermé, même en se limitant à des dynamiques markoviennes décrites par une équation maîtresse. En effet, la résolution numérique de l'équation maîtresse requière de nombreuses multiplications matrice-matrice comportant chacune d^2 éléments, où $d = 2^n$ pour n qubits, alors que la résolution de l'équation de Schrödinger se limite à des multiplications matrice-vecteur avec des vecteurs d'état comportant d éléments. L'utilisation de matrices densité et de potentiellement plusieurs opérateurs de bruit ajoute des contraintes significatives de mémoire numérique. De plus, l'optimisation de ces dynamiques nécessite typiquement des centaines d'itérations, même en utilisant une approche par descente de gradient qui est typiquement beaucoup plus efficace qu'une approche sans gradient pour des problèmes possédant plusieurs paramètres (> 10) [188, 189]. Ces gradients, les dérivées premières d'une fonction de coût arbitraire par rapport aux paramètres libres du problème, sont couramment calculés en résolvant l'équation maîtresse en sens inverse du temps, c'est-à-dire en intégrant de $t = T > 0$ à $t = 0$ [106]. La décohérence génère alors des dynamiques qui se dilatent au lieu de se contracter dans le temps, ce qui rend l'équation différentielle raide et peut mener à des divergences numériques.

On comprend alors qu'une approche précise et efficace numériquement permettant d'évaluer les gradients d'une fonction de coût arbitraire basée sur les dynamiques générées par une équation maîtresse est au cœur du contrôle optimal en systèmes quantiques ouverts. C'est précisément ce que nous proposons dans le cadre de ce projet de recherche.

Il existe diverses approches dans la littérature afin de calculer ces gradients pour des systèmes quantiques ouverts, mais elles possèdent des limitations importantes. D'un côté, les approches basées sur le calcul analytique des gradients [111, 190, 191] ne sont pas généralisables à des fonctions de coût arbitraires, par exemple lors de l'utilisation de fonctions non analytiques, et sont fastidieuses à utiliser alors que tout ajout à la fonction de coût nécessite un travail de formulation mathématique. De façon plus importante encore, ces approches reposent sur diverses approximations des dynamiques afin d'obtenir une forme analytique des gradients, telles qu'une décomposition de Trotter-Suzuki ou une extension en série de Magnus [192, 193]. Ces approximations affectent l'exactitude des gradients et des dynamiques simulées, ce qui peut grandement affecter la convergence ainsi que la pertinence expérimentale des solutions obtenues.

Une approche prometteuse pour surmonter ces limitations consiste à utiliser une approche purement numériquement pour calculer les gradients d'une transformation arbitraire : la différentiation automatique [157]. Cette approche accumule un graphe de calcul de toutes les opérations numériques réalisées impliquant les paramètres voulus ainsi que leurs dérivées premières. Par la suite, une simple application de la règle de dérivation en

chaîne permet d'obtenir tous les gradients recherchés. Pour une introduction plus complète à la différentiation automatique dans le contexte du contrôle optimal, je réfère la lectrice et le lecteur intéressés aux références suivantes : [194, 195]. Pour nos besoins, l'idée principale est que la différentiation automatique est très utile pour calculer des gradients arbitraires, mais elle vient avec un coût élevé en mémoire numérique afin de construire et conserver le graphe de calcul. Par exemple, pour une transformation $y = x^2$ impliquant le paramètre x , la dérivée première est $dy/dx = 2x$. Alors que le calcul numérique se poursuit avec d'autres transformations, à la fois le résultat y et la dérivée $2x$ doivent être conservés en mémoire pour l'évaluation ultérieure des gradients. Dans le contexte du contrôle optimal en système ouvert, les transformations réalisées sont associées à l'intégration de l'équation maîtresse. Ainsi, chaque pas d'intégration numérique $\rho(t + \Delta t) = \mathcal{L}\rho(t)$ nécessite de conserver en mémoire à la fois $\rho(t)$ et $\rho(t + \Delta t)$. Pour un grand espace d'Hilbert, ce besoin en mémoire devient simplement insurmontable en pratique et nous avons besoin d'une approche alternative.

6.2 Résultats principaux

Les principales contributions associées à ce projet de recherche sont l'introduction d'une nouvelle méthodologie ainsi que son utilisation dans la démonstration d'un avantage marqué à utiliser le contrôle optimal quantique afin de mesurer et d'initialiser des qubits supraconducteurs. Dans l'article qui suit, nous introduisons ainsi une méthodologie générale et efficace afin de résoudre des problèmes arbitraires de contrôle optimal dans des systèmes quantiques ouverts complexes. Cette tâche était très coûteuse, voire impossible à réaliser avec les approches existantes. Par la suite, nous démontrons des améliorations d'un facteur deux, à la fois sur la fidélité et la durée des opérations, dans la mesure dispersive et la réinitialisation d'un transmon. Ces tâches représentent deux exemples de problèmes d'optimisation très difficiles à résoudre numériquement, lesquels deviennent accessibles grâce à notre approche.

6.2.1 Notre méthodologie

Inspiré par l'article pionnier en apprentissage automatique inspiré par la physique de Chen et co-auteurs [196], nous formulons une approche pour obtenir, avec un faible coût en mémoire numérique, les gradients d'une fonction de coût arbitraire dépendant des dynamiques produites par une équation maîtresse sous forme de Lindblad. Il est important de souligner que cet article largement cité de même que notre approche s'appuient fortement sur la méthode de l'état adjoint introduite dans le domaine mathématique du contrôle il y a plus de soixante ans [106]. Combinant cette méthode de l'état adjoint aux approches

modernes de rétropropagation des gradients [197], nous sommes en mesure à la fois (1) de généraliser notre approche à des problèmes de contrôle arbitraires, c'est-à-dire avec des objectifs et des contraintes complètement généralisables, ainsi que (2) de minimiser les coûts en mémoire de sorte à pouvoir utiliser des processeurs graphiques (GPUs) et rendre pratique l'ingénierie des contrôles appliqués sur des systèmes quantiques ouverts possédant de grands espaces d'Hilbert. Ces aspects rendent notre approche particulièrement prometteuse dans la résolution de divers problèmes de contrôle optimal quantique, ainsi que toutes autres applications bénéficiant de l'utilisation de gradients telles que l'estimation de paramètres, la tomographie d'états et de processus quantiques, ou la spectroscopie du bruit quantique.

L'idée au cœur de notre approche, la méthode de l'état adjoint, consiste d'abord à définir une représentation duale de la matrice densité $\phi(t)$, laquelle décrit la sensibilité de la fonction de coût par rapport à l'état quantique évolué (la matrice densité au temps t). Par la suite, cet état adjoint est évolué dans le sens inverse du temps à partir de l'équation différentielle originale (l'équation maîtresse de Lindblad) afin d'obtenir tous les gradients de la fonction de coût. Voyons plus en détail comment cette approche fonctionne.

Dans un problème de contrôle optimal quantique générique, nous cherchons à obtenir les paramètres de contrôle $\theta = (\theta_1, \dots, \theta_m)$ réalisant les dynamiques voulues sur le système quantique. Le succès de cette tâche est quantifié par une fonction de coût $\mathcal{C} = \mathcal{C}(\theta, \rho(t_0), \dots, \rho(t_n))$, laquelle dépend des paramètres θ et potentiellement de façon explicite des états quantiques à certains moments intermédiaires de l'évolution $\rho(t_j)$. Une telle dépendance est très utile afin de trouver les contrôles optimaux réalisant l'opération voulue tout en respectant certaines caractéristiques désirables, telles que minimiser la population dans certains niveaux. Par exemple dans les optimisations de la mesure du transmon qui suivent, nous utilisons un coût pénalisant la population de la cavité de mesure au-delà d'un certain seuil $\bar{n}_{\text{crit}}$ pour l'ensemble de l'évolution temporelle afin d'éviter les effets indésirables de transitions vers des états excités du transmon, un phénomène appelé l'ionisation [198].

Afin de minimiser la fonction de coût $\mathcal{C}$, nous cherchons donc à calculer les gradients

$$\frac{d\mathcal{C}}{d\theta} = \frac{\partial \mathcal{C}}{\partial \theta} + \sum_j \frac{\partial \mathcal{C}}{\partial \rho(t_j)} \frac{\partial \rho(t_j)}{\partial \theta}. \quad (6.1)$$

Les gradients $\partial \mathcal{C} / \partial \theta$ peuvent être évalués directement à partir de la définition de la fonction de coût lorsqu'il y a une dépendance explicite sur les paramètres θ , soit de façon analytique ou purement numérique en utilisant l'autodifférentiation. Cependant, les états quantiques du système à différents temps $\rho(t_j)$ apparaissant dans cette équation sont obtenus en résolvant l'équation maîtresse de Lindblad qui, je le rappelle, prend la forme

$$\frac{d\rho}{dt} = \mathcal{L}\rho = -i[H, \rho] + \sum_k \mathcal{D}[L_k]\rho, \quad (6.2)$$

avec H l'Hamiltonien du système, L_k les opérateurs de dissipation et le super-opérateur $\mathcal{D}[L]\rho = L\rho L^\dagger - \{L^\dagger L, \rho\}/2$. Ainsi, l'utilisation d'une approche de différentiation automatique nécessite de différentier à travers l'ensemble des opérations réalisées lors de l'intégration de l'équation maîtresse, ce qui amène un coût en mémoire prohibitif. Afin de contrer ce problème, il est possible d'obtenir les gradients voulus en solutionnant une seconde équation différentielle tout en évitant de conserver en mémoire de nombreux états quantiques. Il est alors possible d'exploiter un compromis entre la mémoire utilisée et la complexité de calcul nécessaire.

En pratique, on définit alors un état adjoint

$$\phi(t_j) \equiv \frac{d\mathcal{C}}{d\rho(t_j)}, \quad (6.3)$$

lequel capture la sensibilité de la fonction de coût par rapport à l'état quantique du système. Comme présenté dans l'article, il est aisé de démontrer que l'évolution de cet état adjoint obéit à une équation différentielle analogue à l'équation maîtresse,

$$\frac{d\phi}{dt} = -\mathcal{L}^\dagger \phi = -i[H, \phi] - \sum_k \mathcal{D}^\dagger[L_k]\phi, \quad (6.4)$$

où $\mathcal{D}^\dagger[L]\phi = L^\dagger \phi L - \{L^\dagger L, \phi\}/2$. L'intégration de cette équation différentielle à l'aide d'un algorithme numérique, tel qu'une méthode de Runge-Kutta [46], se fait ici du temps final t_n jusqu'au temps initial t_0 . Notons que le signe moins devant les opérateurs de dissipation implique que les dynamiques de ϕ se contractent dans le sens inverse du temps, ce qui assure la stabilité numérique de la solution obtenue. L'état initial utilisé pour résoudre cette équation différentielle $\phi(t_n) = d\mathcal{C}/d\rho(t_n)$ est obtenu directement à partir de l'état final $\rho(t_n)$ suite à l'intégration de l'équation maîtresse originale Éq. (6.2).

Somme toute, les gradients de la fonction de coût par rapport à l'ensemble des paramètres du problème peuvent être obtenus directement à partir de la connaissance de la matrice densité et de l'état adjoint du système aux temps voulus,

$$\frac{d\mathcal{C}}{d\theta} = \frac{\partial \mathcal{C}}{\partial \theta} - \int_{t_n}^{t_0} \partial_\theta \text{Tr} \left[\phi^\dagger(t) \mathcal{L}(t, \theta) \rho(t) \right] dt. \quad (6.5)$$

Notre approche procède donc en deux étapes réalisées de façon successive : une intégra-

tion directe de l'équation maîtresse suivie par une intégration en sens inverse du temps de l'équation maîtresse et de l'équation adjointe Éq. (6.4). De ce fait, l'utilisation en mémoire est réduite à seulement $\mathcal{O}(n_c \times d^2)$, avec d la dimension de l'espace d'Hilbert et n_c le nombre de points temporels nécessaires à l'évaluation de la fonction de coût. Pour plusieurs problèmes de contrôle optimal, uniquement l'état final du système est utilisé dans la fonction de coût et $n_c = 1$. De plus, il est possible de calculer le scalaire associé à la fonction de coût partielle directement lorsque le temps $t = t_j$ est atteint lors de l'intégration numérique, de sorte à éviter le besoin de garder l'état $\rho(t_j)$ en mémoire.

Le coût en mémoire est donc réduit à $\mathcal{O}(d^2)$ avec notre approche, une réduction de plusieurs ordres de magnitude pour de nombreux problèmes de contrôle par rapport à une approche par différentiation automatique [194, 199]. En effet, cette dernière approche nécessite plutôt un coût en mémoire de $\mathcal{O}(m \times d^2)$, où m est le nombre d'évaluations de l'équation différentielle utilisé par l'algorithme d'intégration. Par exemple, le stockage de la matrice densité pour un système possédant un espace d'Hilbert de dimension $d = 1000$ et une évolution temporelle avec $m = 2000$ pas d'intégration nécessite environ 30 GB de mémoire, ce qui dépasse déjà les capacités de nombreux processeurs graphiques (GPUs). En revanche, notre approche ouvre la porte à la résolution de problèmes de contrôles dans de grands systèmes quantiques ouverts ($N \lesssim 10\,000$) ainsi que pour des problèmes nécessitant de nombreux pas d'intégration (aucune dépendance en m en principe¹). Je note que typiquement le nombre de pas d'intégration nécessaire croît en fonction du ratio entre l'échelle de temps des dynamiques significatives les plus rapides du système (par exemple la fréquence d'un terme de pilotage micro-onde) et la durée totale d'évolution.

6.2.2 Mesure dispersive optimale d'un transmon

Nous utilisons notre approche afin de trouver les contrôles réalisant une mesure dispersive optimale d'un transmon décrit par un modèle exhaustif incluant à la fois l'Hamiltonien complet du transmon [27] ainsi que sa cavité de mesure et son filtre Purcell [200]. Ce modèle nécessite un espace d'Hilbert d'environ $d = 2000$. Nous démontrons d'abord que l'ingénierie de l'impulsion micro-onde utilisée pour la mesure permet d'améliorer à la fois

1. Tel que décrit dans l'article, la résolution de l'équation maîtresse en sens inverse du temps génère des dynamiques en expansion, lesquelles peuvent mener à des instabilités numériques. Ce problème est cependant complètement résolu en gardant en mémoire la matrice densité à quelques points de contrôle lors de l'évaluation en sens normal du temps de la même équation. En pratique, le nombre de points temporels où la matrice densité doit être conservée en mémoire dépend de l'échelle de temps la plus rapide des processus de dissipation. L'utilisation de ces points de contrôles mène donc à une mise à l'échelle de la mémoire en fonction du nombre de pas d'intégration nécessaire, mais avec un pré-facteur qui est typiquement plusieurs ordres de grandeur plus petit que dans le cas de la différentiation automatique.

la fidélité et la durée totale de la mesure par un facteur d'environ 25 %. Ce gain relativement modeste est limité par l'amplitude finie de l'interaction dispersive entre le transmon et le mode du système cavité-filtre utilisé pour la mesure. Intuitivement, le contrôle de la forme de l'impulsion permet de modifier la trajectoire de ce mode de mesure de façon à maximiser la séparation entre la réponse lorsque le qubit se retrouve dans l'état $|g\rangle$ au lieu de $|e\rangle$. Cependant, des trajectoires arbitraires ne peuvent être réalisées parce que la vitesse et la direction avec laquelle ces trajectoires se séparent sont essentiellement limitées par la forme de l'interaction dispersive

$$\hat{H}_{\text{dispersif}} = \chi \hat{a}^\dagger \hat{a} \hat{\sigma}_z, \quad (6.6)$$

créant des rotations dans l'espace de phase du mode de mesure avec un sens de rotation opposé pour $|g\rangle$ et $|e\rangle$ [13]. Ainsi, le nombre maximal de photons admissibles dans la cavité $\bar{n} = \langle \hat{a}^\dagger \hat{a} \rangle < \bar{n}_{\text{crit}}$, la valeur essentiellement fixe de l'interaction dispersive $\chi \approx g^2/\Delta$ et le taux auquel l'information est évacuée du mode de mesure κ limitent tous la fidélité de mesure réalisable pour une durée donnée [13].

L'ajout d'un contrôle micro-onde directement sur le transmon permet cependant à notre approche d'optimiser la fréquence et l'enveloppe de ce signal afin de découvrir deux stratégies de contrôle significativement plus performantes. Dans la première, le contrôle optimisé sur le transmon possède environ la fréquence de la cavité, ce qui crée de façon effective une interaction longitudinale de la forme

$$\hat{H}_{\text{longitudinal}} = g_z (\hat{a}^\dagger + \hat{a}) \hat{\sigma}_z, \quad (6.7)$$

laquelle engendre plutôt un déplacement linéaire dans l'espace de phase du mode de mesure avec une direction opposée pour $|g\rangle$ et $|e\rangle$ [201, 202], ce qui est optimal [203]. La vitesse de cette séparation est cependant encore une fois limitée par la valeur du décalage dispersif entre les états $|g\rangle$ et $|e\rangle$ du transmon [204].

La seconde stratégie d'optimisation trouve plutôt une fréquence d'impulsion micro-onde sur le transmon correspondant à sa transition $|e\rangle$ - $|f\rangle$. Ce contrôle permet de transférer la population du premier au second état excité du transmon. Cette approche, connue sous le nom de mise à l'étage ou *shelving* [205], est bénéfique parce que la réponse du mode de mesure dans le second état excité diffère davantage de la réponse de l'état fondamental, c'est-à-dire que $\chi_{gf} > \chi_{ge}$, créant ainsi un plus grand signal de mesure [13]. Cette approche permet d'atteindre une fidélité de mesure optimale deux fois plus rapidement que l'approche usuelle, soit en 40 ns au lieu de 80 ns pour les paramètres choisis. L'erreur de mesure associée à cette approche est également deux fois plus petite, soit 1.0×10^{-3} au lieu de 2.1×10^{-3} .

Il est également possible d'interpréter clairement plusieurs caractéristiques des contrôles optimaux trouvés, notamment avec cette approche de mise à l'étage où l'impulsion sur le qubit réalise une porte π d'environ 10 ns entre les états $|e\rangle$ et $|f\rangle$ du transmon avec une enveloppe très similaire à celle communément utilisée [41]. Ce qui est particulièrement impressionnant est que la fidélité de cette porte est de plus de 99 % alors qu'une forte impulsion concomitante sur la cavité, laquelle permet d'amorcer la mesure lors de ces 10 ns, cause un décalage ac-Stark significatif et dépendant du temps sur la fréquence du transmon. La calibration expérimentale d'une telle opération serait donc très ardue sans une telle approche de contrôle optimal en mesure d'optimiser simultanément plusieurs paramètres interdépendants.

6.2.3 Réinitialisation optimale d'un transmon

Nous appliquons enfin notre approche de contrôle optimal quantique en système ouvert sur le problème de réinitialisation du transmon dans le même modèle possédant une cavité et un filtre. Cette opération cherche à amener le transmon dans son état $|g\rangle$ de façon inconditionnelle en utilisant typiquement deux impulsions micro-ondes : une à la fréquence $|e\rangle$ - $|f\rangle$ afin de transférer la population du premier au second état excité et la seconde impulsion à la fréquence $|f0\rangle$ - $|g1\rangle$ afin de transférer cette population dans la cavité possédant un grand taux de dissipation $\kappa \gg \gamma_t$ [206]. Dans l'approche habituelle, ces contrôles sont définis à l'aide de quatre paramètres qui sont calibrés un à la fois puis fixés, soient l'amplitude et la fréquence des deux impulsions possédant une enveloppe plate. En revanche, notre approche de contrôle optimal permet d'optimiser l'ensemble de ces paramètres simultanément, et même d'optimiser beaucoup plus de paramètres en modifiant la forme des impulsions micro-ondes utilisées. Nos résultats montrent qu'une telle optimisation du système quantique ouvert permet d'obtenir une réinitialisation significativement plus performante du transmon, avec des populations résiduelles du transmon (dans les états $|e\rangle$ et $|f\rangle$) de moins de 3×10^{-4} en 200 ns. L'approche de calibration habituelle atteint des populations résiduelles dix fois plus grandes pour la même durée ou trois fois plus grandes en plus de 300 ns.

6.3 Publication [5] : Contrôle optimal de grands systèmes quantiques ouverts

Cette section reproduit la publication scientifique [5] dans laquelle je suis deuxième auteur. J'ai proposé de poursuivre cette problématique de recherche, soit l'optimisation de la mesure et de la réinitialisation du transmon, en adaptant l'approche de [196]. J'ai travaillé de près avec le premier auteur, alors en visite pour quatre mois dans le groupe. Principalement, j'ai contribué à l'élaboration du code de simulation et d'optimisation en PyTorch [159], à l'analyse des résultats et à la rédaction du manuscrit.

Nous avons soumis l'article qui suit au journal *Physical Review Letters* et sommes en attente du retour de la seconde ronde d'arbitrage.

Optimal control in large open quantum systems – the case of transmon readout and reset

Ronan Gautier,^{1,2,3} Élie Genois,¹ and Alexandre Blais^{1,4}

¹*Institut Quantique and Département de Physique,
Université de Sherbrooke, Sherbrooke, QC, Canada*

²*Laboratoire de Physique de l'École Normale Supérieure, Inria, ENS,
Mines ParisTech, Université PSL, Sorbonne Université, Paris, France*

³*Alice & Bob, Paris, France*

⁴*Canadian Institute for Advanced Research, Toronto, ON, Canada*

We present a framework that combines the adjoint state method together with reverse-time back-propagation to solve prohibitively large open-system quantum control problems. Our approach enables the optimization of arbitrary cost functions with fully general controls applied on large open quantum systems described by a Lindblad master equation. It is scalable, computationally efficient, and has a low memory footprint. We apply this framework to optimize two inherently dissipative operations in superconducting qubits which lag behind in terms of fidelity and duration compared to other unitary operations: the dispersive readout and all-microwave reset of a transmon qubit. Our results show that, given a fixed set of system parameters, shaping the control pulses can yield 2x improvements in the fidelity and duration for both of these operations compared to standard strategies. Our approach can readily be applied to optimize quantum controls in a vast range of applications such as reservoir engineering, autonomous quantum error correction, and leakage-reduction units.

Introduction.— Quantum optimal control (QOC) provides a framework to design external controls for realizing arbitrary quantum operations with maximal fidelity and minimal time [1–3], crucial requirements of useful quantum error correction [4]. A common assumption of QOC is that minimizing operation time will also reduce the impact of environmental noise, such that only closed quantum systems need to be considered. However, these approaches are limited by the fact that controlling the system’s coherent dynamics can drastically alter the impact of some noise sources, as exemplified by dynamical decoupling methods [5–7]. Moreover, closed-system approaches cannot extend to inherently dissipative processes such as qubit readout and reset. Consequently, optimally controlling open quantum systems emerges as an important avenue [8–10]. It addresses both the minimization of decoherence in quantum information processing [11–13] and the design of dissipative protocols [14–16], marking a significant step towards comprehensive quantum control in engineered systems.

Over the last decade, several approaches have emerged for open-system QOC. Closed-loop control methods based on feedback engineering [17, 18] or reinforcement learning [19–22] have seen recent success but are difficult to scale to a large number of parameters. Open-loop methods such as gradient-ascent pulse engineering [15, 23], Krotov’s method [24–26] or automatic differentiation [27–30] circumvent this limitation by computing numerical gradients, thereby providing a path to explore the parameter space. However, the control of large open systems remains challenging: the sheer size of the problem prohibits any framework that requires vectorizing the Liouvillian, or storing many density matrices. Two recent papers [31, 32] propose to mitigate this memory overhead

with semi-analytical approaches.

In this Letter, we present a framework enabling the realization of QOC on large open quantum systems with a fully general parametrization over the controls and arbitrary cost functions. Our approach combines the adjoint state method [33–35] with reverse-time back-propagation [36–39] to solve prohibitively large open-system quantum control problems defined in Lindblad form. This approach ensures precise and fast numerical computation of arbitrary gradients with minimal memory usage, making it ideal for GPU acceleration. In the second part of this work, we demonstrate the usefulness of this approach by optimizing two critical operations for the realization of a fault-tolerant quantum computer based on superconducting circuits: dispersive readout [40–42] and all-microwave reset [43, 44] of a transmon qubit [45].

Adjoint state method.— Consider a QOC problem for which we seek to find a set of parameters minimizing a cost function $C(\theta, \hat{\rho}(t_0), \dots, \hat{\rho}(t_n))$. This function, in general, depends on both the problem parameters $\theta = (\theta_1, \dots, \theta_m)$ and on the density matrix of the system at a set of times, $\hat{\rho}(t_i)$. Gradient-based approaches to optimize the control parameters rely on computing the derivative of the cost function with respect to each parameter, $dC/d\theta$. To do so, we apply the adjoint state method [33] to open quantum systems. In this context, the adjoint state is defined as $\hat{\phi}(t) = dC/d\hat{\rho}(t)$, and represents how a change in the density matrix at time t modifies the cost function. For open quantum systems under the usual Born-Markov approximations [46], the evolution of the density matrix is governed by a Lind-

blad master equation ($\hbar = 1$),

$$\frac{d\hat{\rho}}{dt} = \mathcal{L}\hat{\rho} \equiv -i[\hat{H}, \hat{\rho}] + \sum_k \mathcal{D}[\hat{L}_k]\hat{\rho}, \quad (1)$$

where $\hat{H}$ is the system Hamiltonian, $\hat{L}_k$ are jump operators, and $\mathcal{D}[\hat{L}]\hat{\rho} = \hat{L}\hat{\rho}\hat{L}^\dagger - \{\hat{L}^\dagger\hat{L}, \hat{\rho}\}/2$. The adjoint state is then subject to a dual ordinary differential equation [47],

$$\frac{d\hat{\phi}}{dt} = -\mathcal{L}^\dagger\hat{\phi} \equiv -i[\hat{H}, \hat{\phi}] - \sum_k \mathcal{D}^\dagger[\hat{L}_k]\hat{\phi}, \quad (2)$$

where $\mathcal{D}^\dagger[\hat{L}]\hat{\phi} = \hat{L}^\dagger\hat{\phi}\hat{L} - \{\hat{L}^\dagger\hat{L}, \hat{\phi}\}/2$. This equation can be integrated numerically over the time interval of interest $[t_0, t_n]$ with initial condition $\hat{\phi}(t_n) = \partial C / \partial \hat{\rho}(t_n)$, computed analytically if a closed form is available, or directly through automatic differentiation. Notably, the overall minus sign in Eq. (2) ensures numerical stability of the integration by generating contracting dynamics in reverse time. The derivative of the cost function with respect to the problem parameters is given by

$$\frac{dC}{d\theta} = \frac{\partial C}{\partial \theta} - \int_{t_n}^{t_0} \partial_\theta \text{Tr}[\hat{\phi}^\dagger(t) \mathcal{L}(t, \theta) \hat{\rho}(t)] dt. \quad (3)$$

This integral is straightforward to compute using the density matrix and adjoint state at each time $t \in [t_0, t_n]$, as obtained from Eqs. (1) and (2). In particular, the partial derivative with respect to θ can be easily computed from automatic differentiation of the adjoint state equation by noting that

$$\partial_\theta \text{Tr}[\hat{\phi}^\dagger \mathcal{L} \hat{\rho}] = -\text{Tr}[\partial_\theta (d\hat{\phi}/dt)^\dagger \hat{\rho}], \quad (4)$$

which has the form of a vector-Jacobian product.

The QOC optimization is illustrated in Fig. 1 and proceeds in two steps. First, the forward pass consists in using the initial set of parameters (e.g. a sequence of discrete pulses) to numerically integrate the master equation from t_0 to t_n while saving the density matrix at each time t_i of interest. The cost function $C(\theta, \hat{\rho}(t_0), \dots, \hat{\rho}(t_n))$ is then evaluated. To lower the memory footprint, the cost function can also be evaluated on the fly during the forward pass such that only a single density matrix needs to be stored. In a second step, the backward pass, both the master and adjoint equations are simultaneously integrated in reverse time, starting from $t = t_n$. During this process, the integral of Eq. (3) is iteratively evaluated, such as to obtain the entire gradients $dC/d\theta$ once the backpropagation is finished. Having access to the gradients of the cost function, we can now iteratively update the control parameters using standard optimization algorithms [48–50].

We emphasize how each density matrix (blue) is computed twice: once during the forward pass, and once

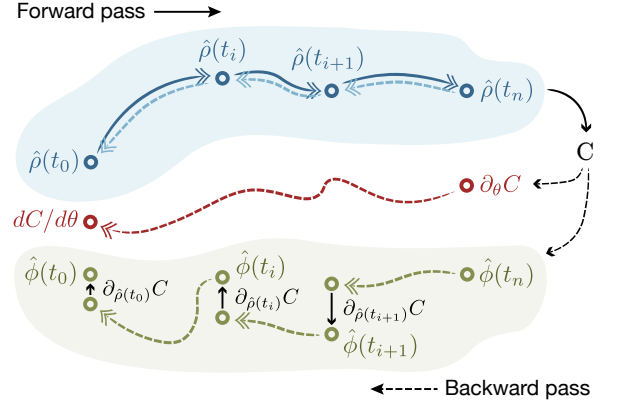


Figure 1. Adjoint state quantum optimal control. In the forward pass, the master equation is integrated and checkpointed for several time points (dark blue). In the backward pass, the density matrix $\hat{\rho}$ is recomputed in reverse time (light blue) together with the adjoint state $\hat{\phi}$ (green) and with the gradients (red) of the cost function C . When a checkpoint is reached, the density matrix is restored to its forward time trajectory, and the adjoint state updated with the corresponding cost function gradient.

during the backward pass. This enables a low memory footprint for the overall scheme, with at most a single density matrix and adjoint state needed to be stored at any given moment. The memory footprint of the method thus scales as $\mathcal{O}(N^2)$ with N the Hilbert space dimension. This is in stark contrast with methods based on automatic differentiation [28], for which the density matrix needs to be stored at each time point of the numerical integration, thus scaling as $\mathcal{O}(M_t N^2)$, with M_t the number of numerical integration steps. Such memory requirements can quickly become prohibitive, even for open quantum systems of intermediate sizes, $N \gtrsim 100$.

Note that this large gain in memory comes at the cost of trading off some numerical runtime. Overall, the scheme requires the integration of four differential equations in total [47], against only two for automatic differentiation. In addition, the reverse time integration of Eq. (1) can be numerically unstable due to the expansive dynamics of the system. This can however be fully resolved with checkpointing the quantum states during the forward pass [39], thus effectively trading back some memory for numerical stability. In practice, checkpointing at the time scale of the largest dissipation operator is sufficient to ensure numerical stability without adding significant complexity.

We have implemented this optimization scheme using PyTorch [51], taking advantage of its automatic differentiation capabilities and GPU support. This framework allows us to run optimization problems for open quantum system with hundreds of parameters, arbitrary cost functions, and for Hilbert space dimensions of up to $N \sim 5000$ while running on a single GPU with 24 GB of memory. Our code is available through the `dynamics` open-source

library [52], simplifying replication of this work and its application to various QOC problems. We now demonstrate the usefulness of this method by optimizing readout and reset of a transmon, two operations that inherently rely on dissipation.

Transmon model.— Let us consider the experimentally realistic model depicted in Fig. 2(a) of a transmon coupled to a readout resonator and Purcell filter [53]

$$\frac{d\hat{\rho}}{dt} = -i[\hat{H}, \hat{\rho}] + \gamma\mathcal{D}[\hat{b}]\hat{\rho} + \kappa\mathcal{D}[\hat{f}]\hat{\rho}, \quad (5)$$

with transmon relaxation rate γ and filter relaxation rate κ , and where

$$\begin{aligned} \hat{H} = & 4E_C\hat{n}_t - E_J\cos(\hat{\varphi}_t) + \omega_r\hat{a}^\dagger\hat{a} + \omega_f\hat{f}^\dagger\hat{f} \\ & - ig\hat{n}_t(\hat{a} - \hat{a}^\dagger) - J(\hat{a} - \hat{a}^\dagger)(\hat{f} - \hat{f}^\dagger) \\ & + \Omega_t\hat{n}_t\sin(\omega_{d,t}t) - i\Omega_f(\hat{f} - \hat{f}^\dagger)\sin(\omega_{d,f}t). \end{aligned} \quad (6)$$

The first two terms denote the free transmon Hamiltonian with charging energy E_C and Josephson energy E_J , with $\hat{n}_t$ and $\hat{\varphi}_t$ the charge and phase operators, and with $\hat{b}$ the corresponding annihilation operator in the diagonal basis. The resonator and filter modes are denoted by $\hat{a}$ and $\hat{f}$, with respective frequencies ω_r and ω_f . These three modes are capacitively coupled in series with coupling strengths $g \gg J$. The system can be driven using a capacitive coupling either through the transmon with a microwave pulse at frequency $\omega_{d,t}$ and envelope $\Omega_t(t)$, or through the Purcell filter at frequency $\omega_{d,f}$ and envelope $\Omega_f(t)$.

For numerical simulation of this model, we first diagonalize the free transmon Hamiltonian and identify the lowest energy eigenstates. We also diagonalize the resonator-filter subsystem yielding two normal modes, each coupled to the transmon. Finally, we apply the rotating-wave approximation (RWA) on couplings and drives. This allows for larger numerical time steps by eliminating fast oscillating dynamics thereby simplifying master equation integration. However, this also implies that not all of the chaotic or transmon ionization dynamics are captured [54–56]. To avoid probing these regimes, we limit the maximum amplitudes of control drives, e.g. to 200 MHz for transmon readout.

We use typical device parameters corresponding to a critical photon number of $\bar{n}_{\text{crit}} = (\Delta/2g)^2 = 16$ [40] and dispersive rates of $\chi/2\pi = 3.8$ MHz and 8.1 MHz with the lower and higher normal modes, respectively. The filter loss rate is $\kappa/2\pi = 30$ MHz and the transmon relaxation time is $T_1 = 20$ μ s. The remaining system parameters can be found in [47].

Transmon readout.— Readout of transmon qubits is realized through the dispersive coupling to a resonator [40]. In this case, the resonator frequency is shifted by the average occupancy in the transmon, and can be measured by driving the resonator at its bare frequency and monitoring the output field, see Fig. 2(b).

In the presence of a Purcell filter, either normal mode of the hybridized resonator-filter subsystem can be used for readout [57].

The metric we use to maximize the measurement fidelity is the signal-to-noise ratio (SNR). Accounting for optimal weighting functions, it reads [58]

$$\text{SNR}(\tau_m) = \sqrt{2\eta\kappa \int_0^{\tau_m} dt |\beta_e(t) - \beta_g(t)|^2}, \quad (7)$$

where $\eta \in [0, 1]$ is the measurement efficiency, τ_m is the readout integration time, and $\beta_{e/g} = \text{Tr}[\hat{f}\hat{\rho}_{g/e}]$ is the average field value in the filter mode, with $\hat{\rho}_{g/e}$ the density matrix obtained after initializing the transmon in the $|g/e\rangle$ state. To obtain results that can be compared to experiments, we use $\eta = 0.6$ [41]. The optimization objective is to maximize the SNR, and thus maximize the distance between the pointer states $|\beta_e - \beta_g|$ in the shortest possible time. Further assuming that the pointer states $\beta_{g,e}$ are Gaussian, one can link the SNR and the transmon lifetime to the readout assignment error [47, 59].

To optimize the transmon readout, we discretize the control pulse envelopes $\Omega(t)$ with 1 ns time bins and use a 250 MHz gaussian filter to interpolate between these pixels during numerical integration and to model realistic experimental distortions [60]. In addition to the discretized drive amplitudes, the optimization parameters θ include the carrier frequency ω_d of each drive. Contrary to the drive amplitudes, the latter are kept constant throughout the pulse duration, in accordance with typical experiments. The cost function used to optimize the transmon readout is principally composed of the SNR of Eq. (7), with additional cost terms constraining the control pulses in order to regularize the optimization and avoid out-of-model dynamics. For example, we limit the number of photons in the hybridized resonator-filter modes, penalize unwanted transitions to higher excited transmon states, and limit the maximal available pulse amplitudes. The full cost function is detailed in [47]. We perform gradient descent using Adam [49] and use the adjoint state method previously described to compute gradients.

Figure 2(c) shows the SNR and the assignment error as a function of the integration time τ_m obtained by our approach, and panel (d) shows the corresponding pulse envelopes for $\tau_m = 40$ ns optimizations. As a point of comparison, we first consider the two non-optimized reference pulses labelled ‘flat’ and ‘two-step’. The former consists of a constant pulse with 2 ns ramp-up and ramp-down times (dark blue squares), and the latter of a two-step pulse meant to rapidly populate the readout mode (blue circles) [41]. In both cases, the amplitude is calibrated to reach $\bar{n} = \bar{n}_{\text{crit}}$ photons in the steady state. The SNR versus τ_m for these two pulses is fitted with the function (full blue and dark blue lines) [53]

$$\text{SNR}(\tau_m) = \alpha\sqrt{2\eta\kappa}(\sqrt{\tau_m} - \sqrt{\tau_{m,0}}), \quad (8)$$

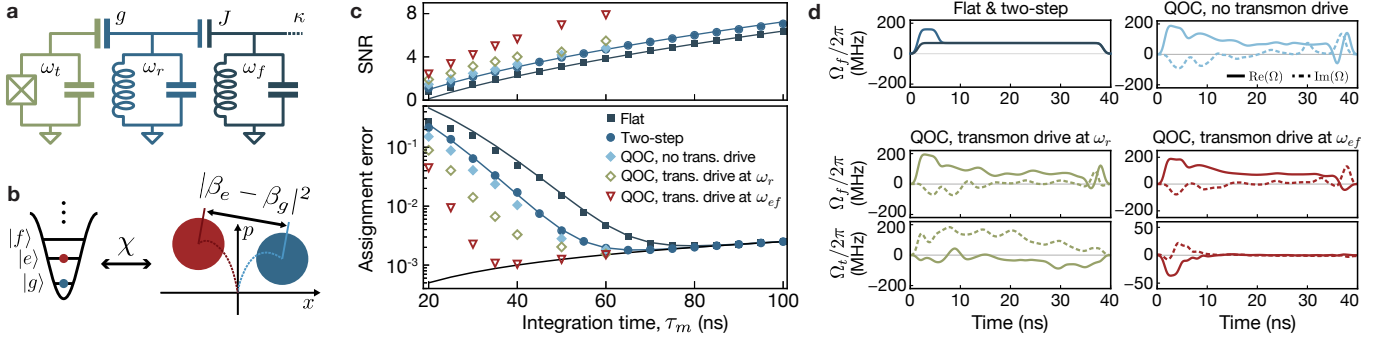


Figure 2. (a) Lumped-element model of a transmon coupled to a readout resonator and Purcell filter. The transmon and filter are driven, and the filter output field is measured through its transmission line. (b) Dispersive readout of a transmon. The mean field in the filter depends on the transmon state. Signal-to-noise ratio of readout increases with the integrated difference of mean fields. (c) Signal-to-noise ratio and assignment error of transmon readout for several drive envelopes: a flat envelope, a 4 ns two-step envelope [41], and optimized envelopes (QOC) with optional additional drives on the transmon at frequencies $\omega_{d,t} \simeq \omega_r$ (green) and $\omega_{d,t} \simeq \omega_{ef}$ (red). The flat and two-step data points are fitted according to Eq. (8). The black line shows the T_1 limit given by $\tau_m/2T_1$. (d) Reference and optimized pulse envelopes at 40 ns of integration time.

where $\alpha = 2|\Omega_f \sin(2\phi)|/\kappa$ is the effective resonator displacement in the steady state, with $\phi = \arctan(2\chi/\kappa)$ and χ the dispersive shift obtained from exact diagonalization of Eq. (6). In this expression, $\sqrt{\tau_{m,0}}$ accounts for an initial delay for the resonator to populate, and is numerically fitted to $\tau_{m,0} = 19$ ns and $\tau_{m,0} = 13$ ns for the flat and two-step pulses, respectively. As the integration time increases, the SNR (assignment error) of both references pulses increase (decrease), up until the transmon T_1 limit is reached (solid black line). Minimum assignment errors of 2.1×10^{-3} and 1.8×10^{-3} are obtained at 80 ns and 65 ns respectively. This is similar performance to state-of-the-art readout experiments [41, 42, 57, 61], as expected from our choice of realistic experimental parameters. Our objective is now to obtain smaller assignment errors in shorter measurement times.

The light blue symbols in Fig. 2(c) are obtained by optimizing the pulse envelope and drive frequency using our QOC approach. The gain is modest and mainly limited by the dispersive coupling with the transmon. Interestingly, the optimized pulses follow a two-step-like shape with a strong initial drive and a weaker subsequent drive, see panel (d). We attribute the small oscillations in the envelope to the rotational gauge freedom of the resonators, which the optimizer is arbitrarily choosing.

Significant improvements are, however, obtained by adding a drive on the transmon concurrently to the readout drive on the resonator. Interestingly, the optimizer converges on two distinct frequencies for the transmon drive. The first strategy found by the optimizer is to drive the transmon at a frequency close to the resonator frequency (green symbols). In that case, the assignment errors decreases faster with integration time than with the above approaches, leading to a minimal assignment error of 1.6×10^{-3} at 60 ns. The effectiveness of this optimized readout strategy stems from the fact that driving the

qubit at the resonator frequency creates a longitudinal-like interaction that can be combined with the usual dispersive interaction to improve readout, as demonstrated in Refs. [62–64].

The second strategy found by the optimizer employs a transmon drive at the (ac-Stark shifted) $|e\rangle$ - $|f\rangle$ transition frequency (red symbols). Given that the cavity response differs more significantly between the transmon states $|g\rangle$ and $|f\rangle$ than between $|g\rangle$ and $|e\rangle$ [53], transferring population into the $|f\rangle$ state leads to a significant improvement of the assignment error, which reaches 1.0×10^{-3} in 40 ns. Interestingly, this shelving approach has already been used to improve readout in circuit QED [65–68]. There, a π -pulse between $|e\rangle$ and $|f\rangle$ is applied to the transmon followed by the measurement drive. In contrast, the optimized strategy found here applies the π -pulse while the cavity is loaded with measurement photons leading to a considerable reduction in the measurement time, see Fig. 2(d). This is possible because the optimizer accounts for the time-dependent ac-Stark shift. The optimized π -pulse features a DRAG-like envelope [69] and achieves a gate fidelity over 99 % in less than 10 ns, even while the readout mode is being strongly driven. Importantly, we note that this approach could achieve significantly higher fidelities by increasing the modest transmon lifetime of 20 μ s used here, as shown by the high SNR in Fig. 2(c).

Transmon reset.— As a second demonstration of the adjoint state method, we consider the optimization of the $f0$ - $g1$ reset of a transmon [43, 44, 70]. This is an all-microwave reset protocol based on a Raman transition between states $|f00\rangle$ and $|g01\rangle$. For the ket $|ijk\rangle$, $|i\rangle$ stands for the qubit state and $|jk\rangle$ the resonator-filter normal modes. Given the large photon loss rate of the filter, the state $|g01\rangle$ quickly decays to $|g00\rangle$, thus ensuring a fast reset of the transmon $|f\rangle$ state. An additional drive

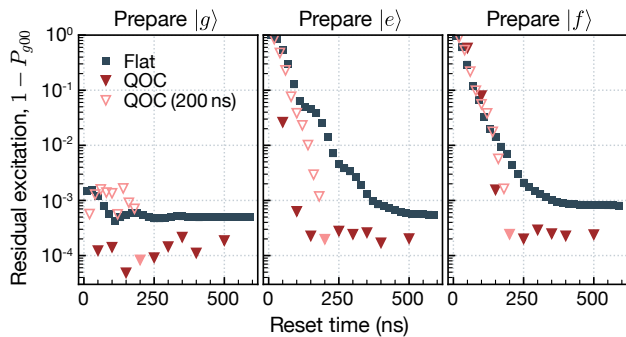


Figure 3. Residual excitation out of $|g00\rangle$ after f0-g1 reset for a calibrated pulse (blue) and an optimized pulse (red) for a system prepared in $|g00\rangle$, $|e00\rangle$, and $|f00\rangle$. Each filled red marker corresponds to a different optimization. Lighter hollow markers illustrate the population dynamics at shorter times for the optimized pulse with 200 ns duration.

at the $|e\rangle$ - $|f\rangle$ transition frequency allows to reset both $|e\rangle$ and $|f\rangle$ states of the transmon. We use the adjoint-state method to find optimal controls for both the f0-g1 and e-f drives simultaneously, in a similar fashion as for optimizing the readout. The cost function is now principally maximizing the transmon population in the $|g00\rangle$ state at the end of the protocol, along with smaller contributions for regularizing the pulses, see [47] for details.

The results of the reset optimization are summarized in Fig. 3. The three panels show the residual excitation out of $|g00\rangle$ against the reset time for a reference flat pulse (blue) and an optimized pulse (red) for different initial transmon states. The reference pulse is composed of two constant drives at the f0-g1 and e-f transitions, where amplitudes and frequencies are calibrated numerically in a similar fashion to what is done in experiments, see [47]. The QOC pulse is obtained by optimizing the carrier frequency and envelopes of both drives, for several total reset times. The optimized pulses show significant improvement over the reference, with a residual excitation of less than 0.05 % at 100 ns (200 ns) for the $|e\rangle$ ($|f\rangle$) state preparation. Note that this delay in the $|f\rangle$ reset time is due to a larger relative weight for the reset of $|e\rangle$ chosen in the cost function, and could be adjusted to achieve the most experimentally relevant reset scheme. This represents a notable improvement over the reference pulses, which reach a steady state after more than 300 ns with larger residual excitations of about 0.07 %. Our results also favourably compare to state-of-the-art experimental realizations of this protocol that reach 1.7 % residual excitations in 100 ns [42], or 0.3 % in 300 ns [44].

Conclusion.— We obtained a fully general framework to optimize open quantum system dynamics in large Hilbert spaces by combining the adjoint state method and reverse-time back-propagation. We have demonstrated the applicability of this method to complex open-system optimization problems using the example of su-

perconducting transmon readout and reset. We stress that our method can readily be applied to optimizing a wide range of quantum control problems where the dissipative dynamics play a significant role such as reservoir (dissipation) engineering [71, 72], autonomous QEC [73, 74], leakage-reduction units [75], quantum cooling, and more. We encourage readers to apply this framework on their own optimal control problems using the open-source library *dynamiqs* [52].

Acknowledgments.— We sincerely thank Alain Sarlette, Ross Shillito, Cristobal Lledó, Pierre Guilmin and Adrien Bocquet for useful discussions. R.G. extends his gratitude to Alain Sarlette for helping organise and support this long-term academic exchange. This work was supported by grant ANR-18-CE47-0005, NSERC, the Canada First Research Excellence Fund, and the Ministère de l'Économie et de l'Innovation du Québec. The numerical simulations were performed using HPC resources at Institut quantique and Inria Paris.

REFERENCES

- [1] A. P. Peirce, M. A. Dahleh, and H. Rabitz, Optimal control of quantum-mechanical systems: Existence, numerical approximation, and applications, *Phys. Rev. A* **37**, 4950 (1988).
- [2] J. Werschnik and E. Gross, Quantum optimal control theory, *Journal of Physics B: Atomic, Molecular and Optical Physics* **40**, R175 (2007).
- [3] C. P. Koch, U. Boscain, T. Calarco, G. Dirr, S. Filipp, S. J. Glaser, R. Kosloff, S. Montangero, T. Schulte-Herbrüggen, D. Sugny, and F. K. Wilhelm, Quantum optimal control in quantum technologies. strategic report on current status, visions and goals for research in europe, *EPJ Quantum Technology* **9**, 10.1140/epjqt/s40507-022-00138-x (2022).
- [4] B. M. Terhal, Quantum error correction for quantum memories, *Rev. Mod. Phys.* **87**, 307 (2015).
- [5] L. Viola, E. Knill, and S. Lloyd, Dynamical decoupling of open quantum systems, *Phys. Rev. Lett.* **82**, 2417 (1999).
- [6] K. Khodjasteh and D. A. Lidar, Fault-tolerant quantum dynamical decoupling, *Phys. Rev. Lett.* **95**, 180501 (2005).
- [7] M. J. Biercuk, H. Uys, A. P. VanDevender, N. Shiga, W. M. Itano, and J. J. Bollinger, Optimized dynamical decoupling in a model quantum memory, *Nature* **458**, 996 (2009).
- [8] R. Schmidt, A. Negretti, J. Ankerhold, T. Calarco, and J. T. Stockburger, Optimal control of open quantum systems: Cooperative effects of driving and dissipation, *Physical review letters* **107**, 130404 (2011).
- [9] T. Schulte-Herbrüggen, A. Spörl, N. Khaneja, and S. Glaser, Optimal control for generating quantum gates in open dissipative systems, *Journal of Physics B: Atomic, Molecular and Optical Physics* **44**, 154013 (2011).
- [10] C. P. Koch, Controlling open quantum systems: tools, achievements, and limitations, *Journal of Physics: Condensed Matter* **28**, 213001 (2016).

- [11] M. H. Goerz, D. M. Reich, and C. P. Koch, Optimal control theory for a unitary operation under dissipative evolution, *New Journal of Physics* **16**, 055012 (2014).
- [12] Z. An, H.-J. Song, Q.-K. He, and D. L. Zhou, Quantum optimal control of multilevel dissipative quantum systems with reinforcement learning, *Phys. Rev. A* **103**, 012404 (2021).
- [13] T. Propson, B. E. Jackson, J. Koch, Z. Manchester, and D. I. Schuster, Robust quantum optimal control with trajectory optimization, *Phys. Rev. Appl.* **17**, 014036 (2022).
- [14] D. J. Egger and F. K. Wilhelm, Optimal control of a quantum measurement, *Phys. Rev. A* **90**, 052331 (2014).
- [15] S. Boutin, C. K. Andersen, J. Venkatraman, A. J. Ferris, and A. Blais, Resonator reset in circuit qed by optimal control for large open quantum systems, *Phys. Rev. A* **96**, 042315 (2017).
- [16] D. Basilewitsch, F. Cosco, N. L. Gullo, M. Möttönen, T. Ala-Nissilä, C. P. Koch, and S. Maniscalco, Reservoir engineering using quantum optimal control for qubit reset, *New Journal of Physics* **21**, 093054 (2019).
- [17] H. Chen, H. Li, F. Motzoi, L. Martin, K. B. Whaley, and M. Sarovar, Quantum proportional-integral (pi) control, *New Journal of Physics* **22**, 113014 (2020).
- [18] R. Porotti, V. Peano, and F. Marquardt, Gradient-ascent pulse engineering with feedback, *PRX Quantum* **4**, 030305 (2023).
- [19] V. V. Sivak, A. Eickbusch, H. Liu, B. Royer, I. Tsioutsios, and M. H. Devoret, Model-free quantum control with reinforcement learning, *Phys. Rev. X* **12**, 011059 (2022).
- [20] V. Sivak, A. Eickbusch, B. Royer, S. Singh, I. Tsioutsios, S. Ganjam, A. Miano, B. Brock, A. Ding, L. Frunzio, *et al.*, Real-time quantum error correction beyond break-even, *Nature* **616**, 50–55 (2023).
- [21] Y. Baum, M. Amico, S. Howell, M. Hush, M. Liuzzi, P. Mundada, T. Merkh, A. R. Carvalho, and M. J. Biercuk, Experimental deep reinforcement learning for error-robust gate-set design on a superconducting quantum computer, *PRX Quantum* **2**, 040324 (2021).
- [22] L. Ding, M. Hays, Y. Sung, B. Kannan, J. An, A. Di Paolo, A. H. Karamlou, T. M. Hazard, K. Azar, D. K. Kim, B. M. Niedzielski, A. Melville, M. E. Schwartz, J. L. Yoder, T. P. Orlando, S. Gustavsson, J. A. Grover, K. Serniak, and W. D. Oliver, High-fidelity, frequency-flexible two-qubit fluxonium gates with a transmon coupler, *Phys. Rev. X* **13**, 031035 (2023).
- [23] N. Khaneja, T. Reiss, C. Kehlet, T. Schulte-Herbrüggen, and S. J. Glaser, Optimal control of coupled spin dynamics: design of nmr pulse sequences by gradient ascent algorithms, *Journal of Magnetic Resonance* **172**, 296 (2005).
- [24] V. Krotov, *Global methods in optimal control theory*, Vol. 195 (CRC Press, 1995).
- [25] M. Goerz, D. Basilewitsch, F. Gago-Encinas, M. G. Krauss, K. P. Horn, D. M. Reich, and C. Koch, Krotov: A python implementation of krotov’s method for quantum optimal control, *SciPost physics* **7**, 080 (2019).
- [26] D. Basilewitsch, C. P. Koch, and D. M. Reich, Quantum optimal control for mixed state squeezing in cavity optomechanics, *Advanced Quantum Technologies* **2**, 1800110 (2019).
- [27] H. Jirari, Optimal control approach to dynamical suppression of decoherence of a qubit, *Europhysics Letters* **87**, 40003 (2009).
- [28] N. Leung, M. Abdelhafez, J. Koch, and D. Schuster, Speedup for quantum optimal control from automatic differentiation based on graphics processing units, *Phys. Rev. A* **95**, 042318 (2017).
- [29] A. G. Baydin, B. A. Pearlmutter, A. A. Radul, and J. M. Siskind, Automatic differentiation in machine learning: a survey, *Journal of Machine Learning Research* **18**, 1 (2018).
- [30] M. Abdelhafez, D. I. Schuster, and J. Koch, Gradient-based optimal control of open quantum systems using quantum trajectories and automatic differentiation, *Phys. Rev. A* **99**, 052327 (2019).
- [31] M. H. Goerz, S. C. Carrasco, and V. S. Malinovsky, Quantum optimal control via semi-automatic differentiation, *Quantum* **6**, 871 (2022).
- [32] Y. Lu, S. Joshi, V. San Dinh, and J. Koch, Optimal control of large quantum systems: assessing memory and runtime performance of grape, *Journal of Physics Communications* **8**, 025002 (2024).
- [33] L. Pontryagin, V. Boltyanskii, R. Gamkrelidze, and E. Mishechenko, *The Mathematical Theory of Optimal Processes* (CRC Press, 1962).
- [34] N. B. Dehaghani and A. P. Aguiar, An application of pontryagin neural networks to solve optimal quantum control problems, *arXiv preprint* (2023), 2302.09143.
- [35] U. Boscain, M. Sigalotti, and D. Sugny, Introduction to the pontryagin maximum principle for quantum optimal control, *PRX Quantum* **2**, 030203 (2021).
- [36] R. T. Chen, Y. Rubanova, J. Bettencourt, and D. K. Duvenaud, Neural ordinary differential equations, *Advances in neural information processing systems* **31**, <https://doi.org/10.48550/arXiv.1806.07366> (2018).
- [37] P. Kidger, *On neural differential equations*, Ph.D. thesis, Mathematical Institute, University of Oxford (2022), 2202.02435.
- [38] J. Somló, V. A. Kazakov, and D. J. Tannor, Controlled dissociation of i2 via optical transitions between the x and b electronic states, *Chemical physics* **172**, 85 (1993).
- [39] S. H. K. Narayanan, T. Propson, M. Bongarti, J. Hückelheim, and P. Hovland, Reducing memory requirements of quantum optimal control, in *International Conference on Computational Science* (Springer, 2022) pp. 129–142.
- [40] A. Blais, R.-S. Huang, A. Wallraff, S. M. Girvin, and R. J. Schoelkopf, Cavity quantum electrodynamics for superconducting electrical circuits: An architecture for quantum computation, *Phys. Rev. A* **69**, 062320 (2004).
- [41] T. Walter, P. Kurpiers, S. Gasparinetti, P. Magnard, A. Potočnik, Y. Salathé, M. Pechal, M. Mondal, M. Oppliger, C. Eichler, and A. Wallraff, Rapid high-fidelity single-shot dispersive readout of superconducting qubits, *Phys. Rev. Appl.* **7**, 054020 (2017).
- [42] Y. Sunada, S. Kono, J. Ilves, S. Tamate, T. Sugiyama, Y. Tabuchi, and Y. Nakamura, Fast readout and reset of a superconducting qubit coupled to a resonator with an intrinsic purcell filter, *Phys. Rev. Appl.* **17**, 044016 (2022).
- [43] M. Pechal, L. Huthmacher, C. Eichler, S. Zeytinoglu, A. A. Abdumalikov, S. Berger, A. Wallraff, and S. Filipp, Microwave-controlled generation of shaped single photons in circuit quantum electrodynamics, *Phys. Rev. X* **4**, 041010 (2014).
- [44] P. Magnard, P. Kurpiers, B. Royer, T. Walter, J.-C. Besse, S. Gasparinetti, M. Pechal, J. Heinsoo, S. Storz,

- A. Blais, and A. Wallraff, Fast and unconditional all-microwave reset of a superconducting qubit, *Phys. Rev. Lett.* **121**, 060502 (2018).
- [45] J. Koch, T. M. Yu, J. Gambetta, A. A. Houck, D. I. Schuster, J. Majer, A. Blais, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf, Charge-insensitive qubit design derived from the cooper pair box, *Phys. Rev. A* **76**, 042319 (2007).
- [46] C. Gardiner and P. Zoller, *Quantum noise: a handbook of Markovian and non-Markovian quantum stochastic methods with applications to quantum optics* (Springer Science & Business Media, 2004).
- [47] See Supplemental Material for more information.
- [48] C. G. Broyden, The convergence of a class of double-rank minimization algorithms 1. general considerations, *IMA Journal of Applied Mathematics* **6**, 76 (1970).
- [49] D. P. Kingma and J. Ba, Adam: A method for stochastic optimization (2014), arXiv:1412.6980.
- [50] H. Robbins and S. Monro, A Stochastic Approximation Method, *The Annals of Mathematical Statistics* **22**, 400 (1951).
- [51] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, T. Killeen, Z. Lin, N. Gimelshein, L. Antiga, *et al.*, Pytorch: An imperative style, high-performance deep learning library, *Advances in neural information processing systems* **32** (2019), 1912.01703.
- [52] P. Guilmin, R. Gautier, A. Bocquet, and E. Genois, dynamiqs: an open-source library for GPU-accelerated and differentiable simulation of quantum systems (2024), in preparation, library available at github.com/dynamiqs/dynamiqs.
- [53] A. Blais, A. L. Grimsmon, S. M. Girvin, and A. Wallraff, Circuit quantum electrodynamics, *Rev. Mod. Phys.* **93**, 025005 (2021).
- [54] J. Cohen, A. Petrescu, R. Shillito, and A. Blais, Reminiscence of classical chaos in driven transmons, *PRX Quantum* **4**, 020312 (2023).
- [55] R. Shillito, A. Petrescu, J. Cohen, J. Beall, M. Hauru, M. Ganahl, A. G. Lewis, G. Vidal, and A. Blais, Dynamics of transmon ionization, *Phys. Rev. Appl.* **18**, 034031 (2022).
- [56] M. F. Dumas, B. Groleau-Paré, A. McDonald, M. H. Muñoz-Arias, C. Lledó, B. D'Anjou, and A. Blais, Unified picture of measurement-induced ionization in the transmon (2024), arXiv:2402.06615.
- [57] F. Swiadek, R. Shillito, P. Magnard, A. Remm, C. Hellings, N. Lacroix, Q. Ficheux, D. C. Zanuz, G. J. Norris, A. Blais, *et al.*, Enhancing dispersive readout of superconducting qubits through dynamic control of the dispersive shift: Experiment and theory (2023), arXiv:2307.07765.
- [58] C. C. Bultink, B. Tarasinski, N. Haandbæk, S. Poletto, N. Haider, D. Michalak, A. Bruno, and L. DiCarlo, General method for extracting the quantum efficiency of dispersive qubit readout in circuit qed, *Applied Physics Letters* **112**, <https://doi.org/10.1063/1.5015954> (2018).
- [59] J. Gambetta, W. A. Braff, A. Wallraff, S. M. Girvin, and R. J. Schoelkopf, Protocols for optimal readout of qubits using a continuous quantum nondemolition measurement, *Phys. Rev. A* **76**, 012325 (2007).
- [60] F. Motzoi, J. M. Gambetta, S. T. Merkel, and F. K. Wilhelm, Optimal control methods for rapidly time-varying hamiltonians, *Phys. Rev. A* **84**, 022307 (2011).
- [61] Y. Sunada, K. Yuki, Z. Wang, T. Miyamura, J. Ilves, K. Matsuura, P. A. Spring, S. Tamate, S. Kono, and Y. Nakamura, Photon-noise-tolerant dispersive readout of a superconducting qubit using a nonlinear purcell filter, *PRX Quantum* **5**, 010307 (2024).
- [62] J. Ikonen, J. Goetz, J. Ilves, A. Keränen, A. M. Gunyho, M. Partanen, K. Y. Tan, D. Hazra, L. Grönberg, V. Vesterinen, S. Simbierowicz, J. Hassel, and M. Möttönen, Qubit measurement by multichannel driving, *Phys. Rev. Lett.* **122**, 080503 (2019).
- [63] S. Touzard, A. Kou, N. E. Frattini, V. V. Sivak, S. Puri, A. Grimm, L. Frunzio, S. Shankar, and M. H. Devoret, Gated conditional displacement readout of superconducting qubits, *Phys. Rev. Lett.* **122**, 080502 (2019).
- [64] M. H. Muñoz Arias, C. Lledó, and A. Blais, Qubit readout enabled by qubit cloaking, *Phys. Rev. Appl.* **20**, 054013 (2023).
- [65] F. Mallet, F. R. Ong, A. Palacios-Laloy, F. Nguyen, P. Bertet, D. Vion, and D. Esteve, Single-shot qubit readout in circuit quantum electrodynamics, *Nat Phys* **5**, 791 (2009).
- [66] B. D'Anjou and W. A. Coish, Enhancing qubit readout through dissipative sub-poissonian dynamics, *Phys. Rev. A* **96**, 052321 (2017).
- [67] S. S. Elder, C. S. Wang, P. Reinhold, C. T. Hann, K. S. Chou, B. J. Lester, S. Rosenblum, L. Frunzio, L. Jiang, and R. J. Schoelkopf, High-fidelity measurement of qubits encoded in multilevel superconducting circuits, *Phys. Rev. X* **10**, 011001 (2020).
- [68] L. Chen, H.-X. Li, Y. Lu, C. W. Warren, C. J. Križan, S. Kosen, M. Rommel, S. Ahmed, A. Osman, J. Biznárová, *et al.*, Transmon qubit readout fidelity at the threshold for quantum error correction without a quantum-limited amplifier, *npj Quantum Information* **9**, 26 (2023).
- [69] F. Motzoi, J. M. Gambetta, P. Rebentrost, and F. K. Wilhelm, Simple pulses for elimination of leakage in weakly nonlinear qubits, *Phys. Rev. Lett.* **103**, 110501 (2009).
- [70] S. Zeytinoglu, M. Pechal, S. Berger, A. A. Abdumalikov, A. Wallraff, and S. Filipp, Microwave-induced amplitude- and phase-tunable qubit-resonator coupling in circuit quantum electrodynamics, *Phys. Rev. A* **91**, 043846 (2015).
- [71] J. F. Poyatos, J. I. Cirac, and P. Zoller, Quantum reservoir engineering with laser cooled trapped ions, *Phys. Rev. Lett.* **77**, 4728 (1996).
- [72] P. M. Harrington, E. Mueller, and K. Murch, Engineered dissipation for quantum information science (2022), arXiv:2202.05280.
- [73] B. Royer, S. Singh, and S. M. Girvin, Stabilization of finite-energy gottesman-kitaev-preskill states, *Phys. Rev. Lett.* **125**, 260509 (2020).
- [74] J. M. Gertler, B. Baker, J. Li, S. Shirol, J. Koch, and C. Wang, Protecting a bosonic qubit with autonomous quantum error correction, *Nature* **590**, 243–248 (2021).
- [75] P. Aliferis and B. M. Terhal, Fault-tolerant quantum computation for local leakage faults, *Quantum Info. Comput.* **7**, 139–156 (2007).
- [76] GPUs used: Nvidia RTX6000 24GB, Nvidia RTX8000 48GB, Nvidia V100 32GB, Nvidia RTX3080 10GB.
- [77] C. Le Bris and P. Rouchon, Low-rank numerical approximations for high-dimensional lindblad equations, *Phys. Rev. A* **87**, 022125 (2013).

- [78] J. R. Dormand and P. J. Prince, A family of embedded runge-kutta formulae, *Journal of computational and applied mathematics* **6**, 19 (1980).

Supplemental Material for "Optimal control in large open quantum systems – the case of transmon readout and reset"

Ronan Gautier,^{1,2,3} Élie Genois,¹ and Alexandre Blais^{1,4}

¹*Institut Quantique and Département de Physique,
Université de Sherbrooke, Sherbrooke, QC, Canada*

²*Laboratoire de Physique de l'École Normale Supérieure, Inria, ENS,
Mines ParisTech, Université PSL, Sorbonne Université, Paris, France*

³*Alice & Bob, Paris, France*

⁴*Canadian Institute for Advanced Research, Toronto, ON, Canada*

CONTENTS

S1. Adjoint state method for open quantum systems	1
A. Adjoint state master equation	1
B. Explicit expression of gradients	1
C. Practical implementation	2
S2. Optimal control of a transmon	2
A. System parameters	2
B. Optimization process	3
C. Numerical integration scheme	3
D. Gaussian filtering	4
E. Cost function contributions	4
S3. Transmon readout	5
A. Assignment error	5
B. Phase-space trajectories	5
S4. Transmon reset	5
A. Calibration	5
B. Pulses	6

Using the definition of the adjoint state and applying the chain rule, we find

$$\begin{aligned}\hat{\phi}(t) &= \frac{dC}{d\hat{\rho}(t)} = \frac{dC}{d\hat{\rho}(t+\varepsilon)} \frac{d\hat{\rho}(t+\varepsilon)}{d\hat{\rho}(t)} \\ &= \hat{\phi}(t+\varepsilon) (1 + \varepsilon \mathcal{L}(t, \theta) + \mathcal{O}(\varepsilon^2)).\end{aligned}\quad (\text{S3})$$

Then, the proof of the adjoint state master equation follows directly from the definition of the derivative,

$$\begin{aligned}\frac{d\hat{\phi}(t)}{dt} &= \lim_{\varepsilon \rightarrow 0} \frac{\hat{\phi}(t+\varepsilon) - \hat{\phi}(t)}{\varepsilon} \\ &= - \lim_{\varepsilon \rightarrow 0} \left(\hat{\phi}(t+\varepsilon) \mathcal{L}(t, \theta) + \mathcal{O}(\varepsilon^2) \right) \\ &= -\hat{\phi}(t) \mathcal{L}(t, \theta)\end{aligned}\quad (\text{S4})$$

or equivalently, using the fact that the adjoint state is hermitian, $\hat{\phi}^\dagger = dC/d\hat{\rho}^\dagger = dC/d\hat{\rho} = \hat{\phi}$, we get the following result thus completing the proof

$$\frac{d\hat{\phi}(t)}{dt} = -\mathcal{L}^\dagger(t, \theta) \hat{\phi}(t). \quad (\text{S5})$$

S1. ADJOINT STATE METHOD FOR OPEN QUANTUM SYSTEMS

A. Adjoint state master equation

In this section, we prove that the adjoint state follows the differential equation Eq. (2) shown in the main text, following a similar derivation as in Ref. [36]. Starting from a general Lindblad master equation

$$\frac{d\hat{\rho}}{dt} = \mathcal{L}\hat{\rho} \equiv -i[\hat{H}, \hat{\rho}] + \sum_k \mathcal{D}[\hat{L}_k] \hat{\rho}, \quad (\text{S1})$$

we can write that for an ε change in time, the evolved quantum state takes the form

$$\begin{aligned}\hat{\rho}(t+\varepsilon) &= \hat{\rho}(t) + \int_t^{t+\varepsilon} \mathcal{L}(\tau, \theta) \hat{\rho}(\tau) d\tau \\ &= \hat{\rho}(t) + \varepsilon \mathcal{L}(t, \theta) \hat{\rho}(t) + \mathcal{O}(\varepsilon^2).\end{aligned}\quad (\text{S2})$$

B. Explicit expression of gradients

To derive an explicit expression for $dC/d\theta$, we also follow Ref. [36] and introduce an augmented density matrix $\hat{\rho}_{aug}(t) = [\hat{\rho}(t), \theta(t)]^T$ which includes a second block row with time-dependent parameters. In practice, parameters are constant throughout the integration, but this fictitious time-dependence will allow us to isolate the sensitivity of the cost function to an infinitesimal change of parameters at each point in time. The augmented density matrix follows the differential equation

$$\frac{d\hat{\rho}_{aug}(t)}{dt} = \mathcal{L}_{aug}(t, \theta) \hat{\rho}_{aug}(t) = \begin{pmatrix} \mathcal{L}(t, \theta) \hat{\rho}(t) \\ 0 \end{pmatrix}. \quad (\text{S6})$$

We can also introduce the augmented adjoint state, $\hat{\phi}_{aug}(t) = [\hat{\phi}(t), dC/d\theta(t)]$, where the parameter sensitivity $dC/d\theta(t)$ represents the cost function gradient with respect to the parameters at a given time. We then follow a similar derivation as in the previous section. The

main difference comes in the application of the chain rule of Eq. (S3) which now yields,

$$\begin{aligned}\hat{\phi}_{aug}(t) &= \frac{dC}{d\hat{\rho}_{aug}(t)} = \frac{dC}{d\hat{\rho}_{aug}(t+\varepsilon)} \frac{d\hat{\rho}_{aug}(t+\varepsilon)}{d\hat{\rho}_{aug}(t)} \\ &= \hat{\phi}_{aug}(t+\varepsilon) (1 + \varepsilon \mathcal{M}(t, \theta) + \mathcal{O}(\varepsilon^2))\end{aligned}\quad (\text{S7})$$

where

$$\begin{aligned}\mathcal{M}(t, \theta) &= \frac{d(\mathcal{L}_{aug}(t, \theta) \hat{\rho}_{aug}(t))}{d\hat{\rho}_{aug}(t)} \\ &= \begin{pmatrix} \mathcal{L}(t, \theta) & \partial_\theta \mathcal{L}(t, \theta) \hat{\rho}(t) \\ 0 & 0 \end{pmatrix},\end{aligned}\quad (\text{S8})$$

and where ∂_θ denotes the partial derivative with respect to θ . Similarly as before, we can rearrange this expression, reinsert it in the definition of the derivative and take the limit of $\varepsilon \rightarrow 0$. Then, selecting only the second block of the resulting differential equation yields

$$\frac{d}{dt} \left(\frac{dC}{d\hat{\rho}(t)} \right) = -\partial_\theta \text{Tr} \left[\hat{\phi}^\dagger(t) \mathcal{L}(t, \theta) \hat{\rho}(t) \right]. \quad (\text{S9})$$

We have thus obtained a straightforward differential equation on the parameter sensitivity. Integrating this equation from $t = 0$ to T , with the initial condition that $dC/d\theta(T) = \partial C/\partial\theta$ yields

$$\frac{dC}{d\theta} = \frac{\partial C}{\partial\theta} - \int_T^0 \partial_\theta \text{Tr} \left[\left(\mathcal{L}^\dagger(t, \theta) \hat{\phi}(t) \right) \hat{\rho}(t) \right] dt, \quad (\text{S10})$$

which is Eq. (3) in the main text. We can also get the derivative of the loss function with respect to the integration time T from a simple application of the chain rule

$$\frac{dC}{dT} = \frac{dC}{d\hat{\rho}(T)} \frac{d\hat{\rho}(T)}{dT} = \text{Tr} \left[\hat{\phi}^\dagger(T) \mathcal{L}(T, \theta) \hat{\rho}(T) \right]. \quad (\text{S11})$$

Finally, between Eqs. (S5), (S10) and (S11), we obtain the gradients of the loss function with respect to all relevant objects involved in the computation, thus allowing one to perform gradient descent for the desired parameters from arbitrary cost functions.

C. Practical implementation

The full implementation in Python of the approach described in the main text is available on the freely available open-source software dynamiqs [52]. To summarize the structure of a numerical implementation of the adjoint state method for open quantum systems, we present a pseudo-code implementation in Alg. S1. It involves a function `ODEstep(f, t, y, θ)` that computes a single step of the integration of a linear ordinary differential equation, $\dot{y} = f(t, \theta)y$, at time t . The `CostFn(θ, Lρ)` function computes the cost function from the parameters θ and a

Algorithm S1: Adjoint state method for open quantum systems.

```

1  t ← t0;
2  ρ ← ρ0;
3  Lρ ← {ρ0}; // checkpoint storage list (optional)
4  for ti ∈ {t1, ..., tn} do // forward pass
5      while t < ti do
6          ρ, δt ← ODEstep(L, t, ρ, θ);
7          t ← t + δt
8      end
9      Lρ ← Lρ + {ρ};
10 end
11 C ← CostFn(θ, Lρ);
12 φ ← AutoDiff(CostFn, C, Lρ[n]);
13 ∇ ← AutoDiff(CostFn, C, θ);
14 for ti ∈ {tn-1, ..., t0} do // backward pass
15     while t > ti do
16         {ρ, φ}, δt ← ODEstep({-L, L†}, t, {ρ, φ}, θ);
17         δφ ← AutoDiff(ODEstep, φ, θ);
18         ∇ ← ∇ - Tr[δφ · ρ];
19         t ← t - δt;
20     end
21     φ ← φ + AutoDiff(CostFn, C, Lρ[i]);
22     ρ ← Lρ[i];
23 end
24 return ∇
```

list of density matrices L_ρ . It can also be computed on the fly in the *for* loop to avoid storing these matrices. Finally, the `AutoDiff(g, xout, xin)` function computes the partial derivative $\partial x_{out}/\partial x_{in}$ through automatic differentiation of g , where x_{in} and x_{out} are given inputs and outputs of the function g . Note that lines 17 and 18 can be merged in a single computation of a vector-Jacobian product to avoid storage of $\delta\phi$.

S2. OPTIMAL CONTROL OF A TRANSMON

A. System parameters

Unless stated otherwise, we use $E_C/2\pi = 315$ MHz, $E_J/E_C = 51$, corresponding to a bare transmon frequency $\omega_t/2\pi = 6$ GHz and to an anharmonicity $\alpha/2\pi = -349$ MHz. The resonator and filter frequencies are $\omega_r/2\pi = 7.2$ GHz and $\omega_f/2\pi = 7.21$ GHz with couplings $g/2\pi = 150$ MHz and $J/2\pi = 30$ MHz. This yields a transmon-resonator detuning of $\Delta/2\pi = 1.2$ GHz, a critical photon number of $\bar{n}_{\text{crit}} = (\Delta/2g)^2 = 16$ [40] and dispersive rates of $\chi/2\pi = 3.8$ MHz and 8.1 MHz with the lower and higher normal modes, respectively. Finally, the filter loss rate is $\kappa/2\pi = 30$ MHz and the transmon relaxation rate is $\gamma/2\pi = 8$ kHz, i.e. $T_1 = 20$ μ s.

Description	Analytical formula	Weight	Hyperparameters
Signal-to-noise ratio	$1/\sqrt{2\eta\kappa} \int_0^{\tau_m} \beta_e(t) - \beta_g(t) ^2 dt$	1.0	$\eta = 0.6, \kappa/2\pi = 30 \text{ MHz}$
Pulse amplitude (filter)	$\frac{1}{\tau_m} \int_0^{\tau_m} \text{ReLU}(\Omega_f(t) - \Omega_{\max}) dt$	0.1	$\Omega_{\max}/2\pi = 200 \text{ MHz}$
Pulse amplitude (transmon)	$\frac{1}{\tau_m} \int_0^{\tau_m} \text{ReLU}(\Omega_t(t) - \Omega_{\max}) dt$	0.1	$\Omega_{\max}/2\pi = 200 \text{ MHz}$
Forbidden states (transmon)	$\sum_{i=g,e} \sum_{k \geq n}^{N_t-1} \frac{1}{\tau_m} \int_0^{\tau_m} \text{Tr}[[k\rangle\langle k _t \hat{\rho}_i(t)] dt$	1.0	$n = 2 \text{ or } 3, N_t = 5$
Forbidden states (undriven normal)	$\sum_{i=g,e} \sum_{k \geq n}^{N_u-1} \frac{1}{\tau_m} \int_0^{\tau_m} \text{Tr}[[k\rangle\langle k _u \hat{\rho}_i(t)] dt$	5000	$n = 2, N_u = 4$
Critical photon number (resonator)	$\sum_{i=g,e} \frac{1}{\tau_m} \int_0^{\tau_m} \text{ReLU}(\text{Tr}[\hat{a}^\dagger \hat{a} \hat{\rho}_i(t)] - \bar{n}_{\text{crit}}) dt$	0.1	$\bar{n}_{\text{crit}} = 16$
State reset $ g\rangle$	$\log_{10}(1 - \langle g00 \hat{\rho}_g(\tau_m) g00 \rangle ^2)$	0.3	–
State reset $ e\rangle$	$\log_{10}(1 - \langle g00 \hat{\rho}_e(\tau_m) g00 \rangle ^2)$	1.0	–
State reset $ f\rangle$	$\log_{10}(1 - \langle g00 \hat{\rho}_f(\tau_m) g00 \rangle ^2)$	0.3	–
Pulse amplitude (transmon)	$\frac{1}{\tau_m} \int_0^{\tau_m} \text{ReLU}(\Omega_t(t) - \Omega_{\max}) dt$	0.1	$\Omega_{\max}/2\pi = 600 \text{ MHz}$

Table S1. Summary of readout (top) and reset (bottom) cost functions. Each line corresponds to a different contribution to the total cost function, with an overall weight tuned heuristically. All contributions are functions of the problem parameters and/or of the density matrix at certain times. Time integrals are numerically discretized in 1 ns time bins. ReLU denotes a rectified linear unit function such that $\text{ReLU}(x) = 0$ for $x \leq 0$ and $\text{ReLU}(x) = x$ for $x \geq 0$. The states $|k\rangle_t$ and $|k\rangle_u$ denote the k -th Fock state in the transmon and undriven normal mode respectively. Finally, $\hat{\rho}_i$ denotes the density matrix for a transmon initialized in state $|i\rangle_t$.

B. Optimization process

In this work, we optimize transmon readout and reset by simulating the Lindblad master equation of Eq. (1) with the Hamiltonian of the main text which contains three modes: the full transmon Hamiltonian, a readout resonator and a Purcell filter. In our simulations, we first diagonalize the transmon Hamiltonian in the charge basis using 300 charge states and then truncate the subsystem to $N_t = 5$ eigenstates of lowest energy. In parallel, we also diagonalize the resonator-filter subsystem yielding two normal modes – each coupled to the transmon – and keep N_d Fock states for the mode used to readout/reset the transmon (driven mode) and N_u Fock states in the other mode (undriven mode). For readout, we typically take $N_d = 85$ and $N_u = 4$. For reset, we instead use $N_d = N_u = 4$ since neither mode is actively driven.

Note that we intentionally select small Hilbert space sizes to lower the simulation runtime: a typical gradient descent on such transmon readout/reset simulations takes a few days per data point using a single GPU [76]. To avoid probing unphysical simulations with such low Hilbert space sizes, we actively steer the optimizer in two ways. First, we engineer the cost function such as to avoid large populations in boundary Fock states and limit the pulse amplitudes. See Table S1 for detailed cost function contributions. Second, we validate all optimized simulations with a single additional simulation in a larger Hilbert space, using pulses as obtained by the optimization process. For such simulations, we typically take $N_t = 6$, and $N_d = 90$ and $N_u = 8$ for readout, or $N_d = N_u = 8$ for reset. All data points shown in the main text correspond to such validated simulations, and

do not deviate significantly from the results obtained on the reduced Hilbert space simulations.

C. Numerical integration scheme

Regarding the integration scheme of the ordinary differential equation (ODE), we used a second-order Runge-Kutta solver [77] with a fixed integration time step of $\delta t = 0.003 \text{ ns}$. Such a method ensures preservation of the positivity and trace of the density matrix at all times during the integration, contrary to standard ODE integration schemes. This is particularly interesting in the context of the adjoint state method to avoid divergence during the reverse-time integration of the master equation. Later on in this work, we have also used a standard order-4/5 Dormand-Prince integration scheme [78] which provided better runtimes in the forward pass and similar ones in the backward pass, even considering the stiffness of the reverse-time integration. In any case, several thousands of integration steps are typically required in our simulations. This makes the use of standard automatic differentiation prohibitive in computer memory. A simple pen-and-paper estimation of memory with the previously stated Hilbert space size and $M_t = 10000$ time steps yields a memory footprint of $8M_t(N_t N_d N_u)^2 = 215 \text{ GB}$ where 8 is the number of bytes used to store a single-precision complex number.

D. Gaussian filtering

As stated in the main text, all simulated pulse envelopes are first discretized in $\tau_0 = 1$ ns time bins and then gaussian filtered according to Ref. [60]. Each complex amplitude of the discretized pulse then corresponds to an additional parameter in the optimization. Concretely, the time-dependent pulses read

$$\Omega(t) = \mathcal{F}_g \left[\sum_{j=0}^{\lfloor \tau_m/\tau_0 \rfloor - 1} \Omega_j \Pi_j(t, \tau_0) \right] (t) \quad (\text{S12})$$

where

$$\Pi_j(t, \tau_0) \equiv \Theta(t - j\tau_0) - \Theta(t - (j+1)\tau_0) \quad (\text{S13})$$

is the rectangle function, with Θ is the Heaviside unit step function, and where $\mathcal{F}_g$ is a gaussian filter of bandwidth $\omega_B/2\pi = 250$ MHz. Commuting $\mathcal{F}_g$ with the sum of Eq. (S12), we get the simpler expression

$$\Omega(t) = \sum_{j=0}^{\lfloor \tau_m/\tau_0 \rfloor - 1} \Omega_j \zeta_j(t), \quad (\text{S14})$$

where

$$\zeta_j(t) = \frac{1}{2} \left[\text{erf} \left(\omega_0 \frac{t - j\tau_0}{2} \right) - \text{erf} \left(\omega_0 \frac{t - (j+1)\tau_0}{2} \right) \right], \quad (\text{S15})$$

with erf the error function and $\omega_0/2\pi = 425.5$ MHz [60].

E. Cost function contributions

Designing appropriate cost functions is essential to obtain useful quantum optimal control results from numerical optimizations. In this work, we define composite cost functions with a principal component and several smaller regularizing contributions. The cost functions used for transmon readout and reset are summarized in Table S1. Let us describe each cost function separately.

For transmon readout, the main contribution is the inverse of the SNR. We use the inverse (instead of e.g. the negation) such that the optimizer focuses on other cost function contributions increasingly more as the SNR grows large. There are five regularizing cost functions on top of this main component. The first two correspond to a maximum pulse amplitude on the filter and transmon respectively, at $\Omega_{\max}/2\pi = 200$ MHz. These cost functions use a rectified linear unit function, or ReLU, such that $\text{ReLU}(x) = 0$ for $x \leq 0$ and $\text{ReLU}(x) = x$ for $x \geq 0$. As such, no hard limits are set on the pulse amplitudes to allow the optimizer to explore a larger parameter space in case interesting solutions can be found at larger amplitudes. Two other contributions correspond

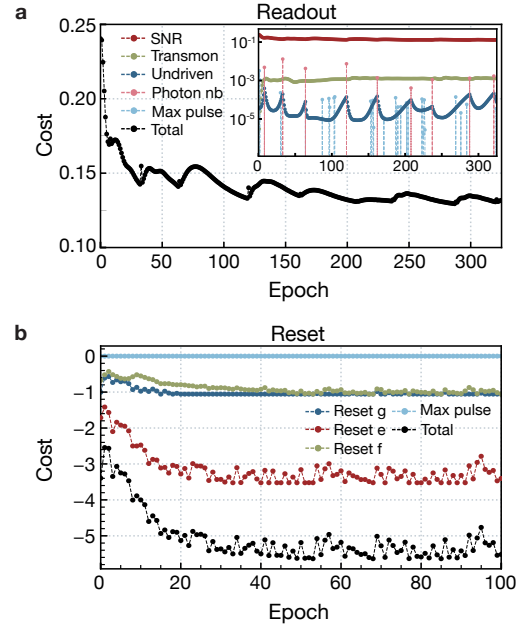


Figure S1. Cost function as a function of the optimization epoch for (a) transmon readout QOC with no transmon drive at 40 ns, and (b) transmon reset QOC at 200 ns. The inset of panel (a) shows the cost function contributions in log scale, with the main contribution being the inversed SNR. Other contributions are meant to regularize the pulse and to avoid unphysical regions. See Table S1 for the detailed cost function contributions.

to forbidden states in the transmon and undriven normal mode respectively. As previously discussed, these costs are used to avoid probing unphysical pulses that would make use of boundary Fock states sitting close to our artificial numerical truncation. Finally, the last contribution to the cost function is a soft upper bound on the number of photons in the resonator. Indeed, according to the standard theory of dispersive qubit readout [53], non-QND readout can be probed by driving the readout resonator above its critical number of photons $\bar{n}_{\text{crit}}$. This cost function also allows for a fair comparison between all readout strategies by limiting the effective speed of pointer state separation, given by $\bar{n}\chi$ in the absence of additional transmon drives.

For transmon reset, the main contributions are the logarithms of residual population outside the $|g00\rangle\langle g00|$ subspace, for a transmon initialized in either $|g\rangle$, $|e\rangle$, or $|f\rangle$. In particular, we weight these three contributions non-uniformly with a larger weight on the $|e\rangle$ state reset. This is a choice motivated by experimental designs for which population inside the qubit subspace is typically much larger than leaked population at the beginning of resets. However, these weights can be freely adapted to the experiment. The only regularizing cost function for transmon reset is an upper bound on the pulse amplitude, similar to that of transmon readout. Because reset attempts to evacuate the system photons, additional cost

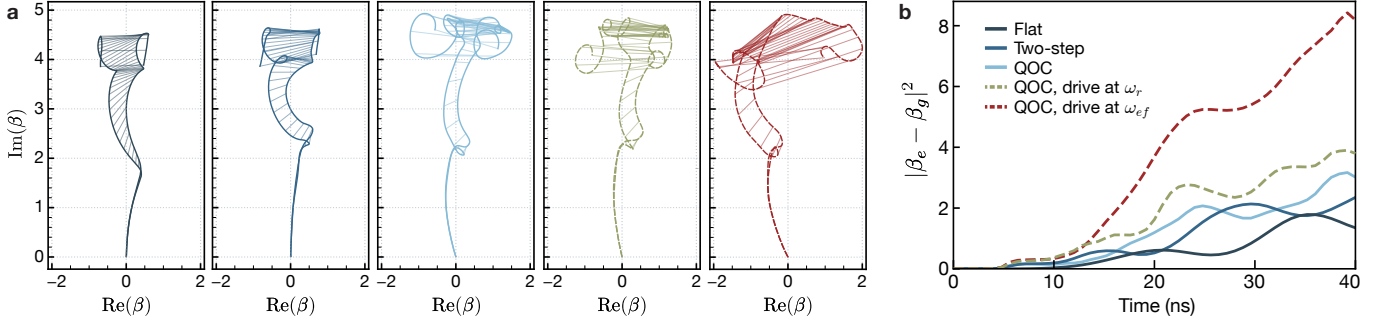


Figure S2. (a) Mean phase-space trajectories of the Purcell filter during a 40 ns readout. Each panel shows β_e and β_g against time for a different readout strategy, where $\beta_{e/g} = \text{Tr}[\hat{f}\hat{\rho}_{g/e}]$ is the filter mean field after initializing the transmon in the $|e\rangle$ or $|g\rangle$ state. Lines that join each trajectory represent 1 ns intervals. The optimal control algorithm attempts to maximize SNR, which corresponds to the integrated distance between both trajectories. (b) Phase-space distance against time for the same readout strategies as in (a).

functions are not required.

In Fig. S1, we show how typical cost function contributions evolve as a function of the optimization epoch (i.e. iteration). Panel (a) shows a 40 ns transmon readout optimization where only the filter is driven, while panel (b) shows a 200 ns reset optimization. In both cases, we find a clear decrease in the total cost in the first 20 to 50 epochs, and convergence to a steady value afterwards. For readout, the cost function is dominated by the SNR cost function. The inset highlights the smaller cost contributions in log scale. For these regularizing cost functions, we regularly see peaks indicating that the optimizer is generating dynamics close to the Hilbert space boundaries before retracting away, all the while achieving increasingly better SNR.

S3. TRANSMON READOUT

A. Assignment error

In Fig. 2(c) of the main text, we show the assignment error of transmon readout as a function of the integration time. This assignment error is computed in Refs. [57, 59] assuming that the transmon pointer states stay gaussian during readout. This assumption is verified to a very good approximation for all of our simulated readout sequences. Given this structure, the assignment error reads

$$\varepsilon_a(\tau_m) \simeq \frac{1}{2} \text{erfc}\left(\frac{\text{SNR}(\tau_m)}{2}\right) + \frac{\tau_m}{2T_1}, \quad (\text{S16})$$

where $\varepsilon_a = 1 - \mathcal{F} = [P(e|g) + P(g|e)]/2$, with $P(i|j)$ the probability of measuring state $|i\rangle$ when $|j\rangle$ is prepared, $\mathcal{F}$ the readout fidelity, and erfc the complementary error function.

B. Phase-space trajectories

Figure S2(a) shows the readout trajectories obtained by our three optimized control strategies, in comparison with the flat and two-step readout references. Notably, we see that pulse shaping the microwave drives on the cavity and qubit, while still using an experimentally-realistic bandwidth, makes the readout dynamics deviate from the usual dispersive trajectories. As quantified in Fig. S2(b), optimizing the control drives leads to both a faster and a larger separation between the qubit pointer states, which directly translates into better readout SNR. The red curve, which corresponds to adding the transmon drive at the $|e\rangle$ - $|f\rangle$ transition frequency such as to achieve shelving, shows how beneficial that approach is for the system parameters considered here, compared to standard dispersive readout. The discovered scheme of simultaneously shelving and populating the readout mode with photons could straightforwardly be implemented with the typical controls available in current transmon experiments.

S4. TRANSMON RESET

A. Calibration

Since the f0-g1 reset of the transmon $|e\rangle$ and $|f\rangle$ states requires two concurrent microwave drives on the transmon, the frequency and amplitude of these drives need to be calibrated in a self-consistent way. In order to get the best reset performance using flat drives on which to compare our optimal controls, we perform a numerical calibration procedure that follows closely what is done in typical experiments [43, 44].

Namely, for a given f0-g1 drive amplitude Ω_{f0g1} , we first scan the reset drive frequency ω_{f0g1} to find the ac-Stark shifted resonant frequency that maps the prepared

$|f\rangle$ transmon state back to $|g\rangle$ after 100 ns. The results of this simulated experiment are illustrated in Fig. S3(a, e) for the lower and higher normal readout mode, respectively, and for different drive amplitudes $\Omega_{f0g1}/2\pi$ between 50 and 650 MHz. Once the frequency ω_{f0g1} is fixed on resonance, we calibrate the $|e\rangle$ - $|f\rangle$ drive frequency ω_{ef} . This is achieved by performing a similar experiment as before, but in which the transmon is prepared in $|e\rangle$ instead of $|f\rangle$, and both the f0-g1 drive and e-f drives are on. The amplitude ratio $\Omega_{f0g1}/\Omega_{ef}$ between both drives is kept constant such that the effective f0-g1 Rabi rate yields $\tilde{g}(\Omega_{f0g1}) = \Omega_{ef}$ [44, 70]. The results of this second calibration are shown in Fig. S3(b, f) for the lower and higher modes, respectively. Finally, the $|e\rangle$ - $|f\rangle$ drive amplitude Ω_{ef} is calibrated from a 1D parameter sweep while performing the same $|e\rangle$ reset experiment. This sweep is shown in Fig. S3(c, g). In the end, this whole procedure yields a single set of calibrated parameters (ω_{f0g1} , ω_{ef} , Ω_{ef}) for every value of Ω_{f0g1} and for each normal mode. We simulate the corresponding reset of each set of calibrated parameters, and show the corresponding results in the panels of Fig. S3(d, h).

Out of all of the calibrated sets of parameters, we selected the best performing reset (both in terms of speed and fidelity) to compare with our optimal controls in Fig. 3 of the main text. This corresponds to a f0-g1 amplitude of $\Omega_{f0g1}/2\pi = 350$ MHz driving the higher frequency normal mode of the coupled resonator-filter subsystem. This mode has a larger dispersive coupling to the transmon than the lower frequency normal mode. This particular reset is shown with a red title in Fig. S3(h).

B. Pulses

In Fig. S4, we present the optimal reset pulse shapes found using our quantum optimal control approach. The 200 ns pulse of panel (b) produces the transmon population dynamics illustrated with the pink markers in Fig. 3 of the main text. Interestingly, in the 100 ns pulse of panel (a), we see that the numerical optimization yields an initial π -pulse to transfer population from $|e\rangle$ to $|f\rangle$ which quickly decays back to $|g\rangle$ because of the f0-g1 drive. This behaviour can be understood by the fact that our chosen cost function puts more weight into resetting the $|e\rangle$ state. As such, the optimal drive scheme with this constraint initially resets the $|e\rangle$ state, and subsequently resets the population that was initially in the $|f\rangle$ state at the beginning of the protocol. This behaviour is consistent with the fact that the residual population when initializing in $|e\rangle$ converges after 100 ns whereas it converges after 200 ns for the $|f\rangle$ initial state, as reported in the main text. Our choice of relative weights between the different residual populations is motivated by the thermal population of a high frequency transmon (5 GHz) in a relatively cold environment (20 mK $\sim$ 417 MHz), but

is arbitrary and could easily be adapted to other experimentally relevant scenarios, e.g. a low-frequency fluxonium qubit (<1 GHz) in contact with the same thermal bath.

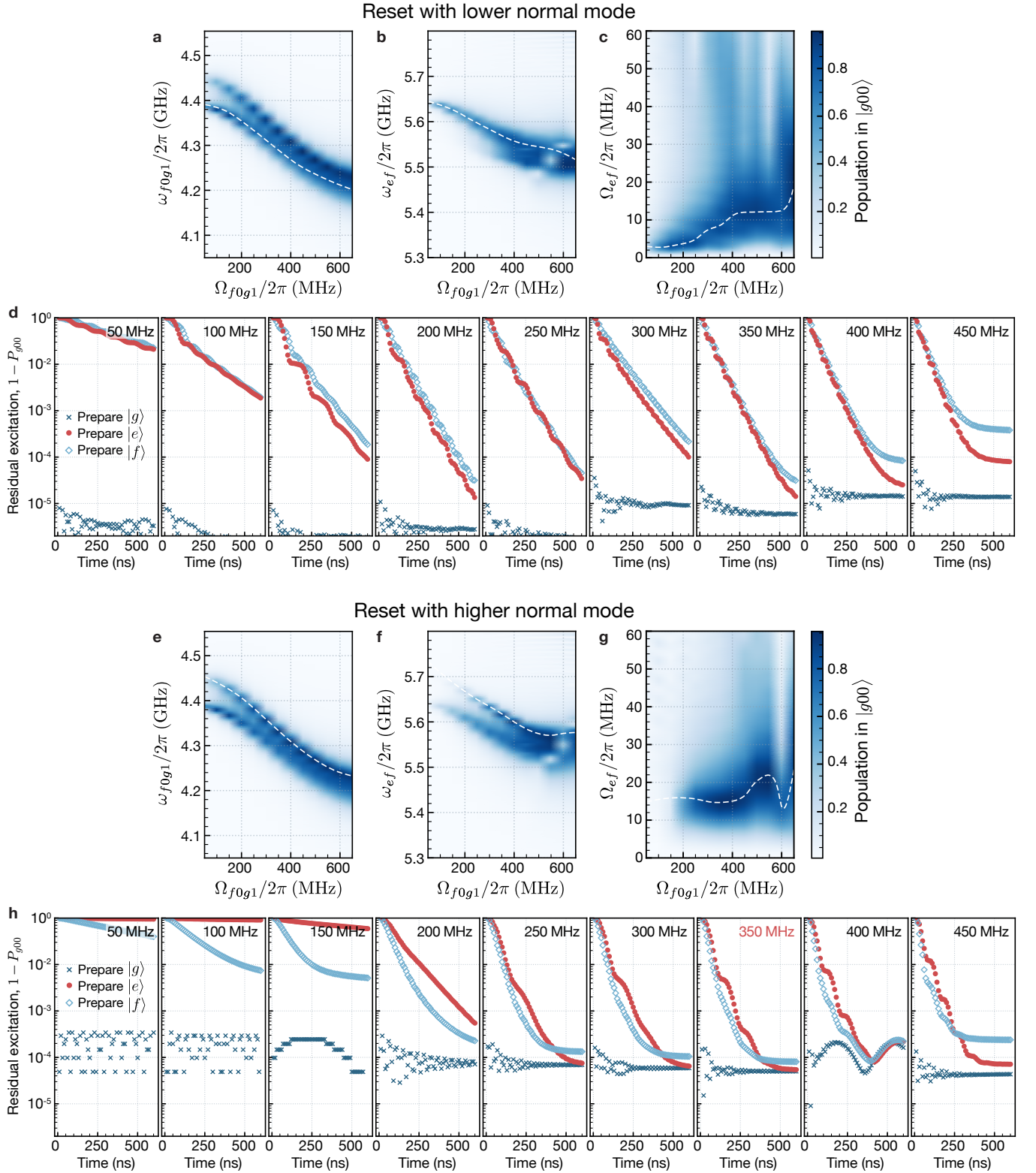


Figure S3. Numerical calibration of the f0-g1 reset for the reference flat drives. (a-c) Following the normal mode of lowest frequency, the f0-g1 drive frequency $\omega_{f0g1}/2\pi$, the e-f drive frequency $\omega_{ef}/2\pi$, and e-f drive amplitude $\Omega_{ef}/2\pi$ are sequentially calibrated. See the protocol description in Sec. S4 for more details. (d) Resulting f0-g1 reset performance after preparing the $|g\rangle$, $|e\rangle$, and $|f\rangle$ states of the transmon for different f0-g1 drive amplitudes $\Omega_{f0g1}/2\pi$ between 50 MHz and 450 MHz. (e-h) Results for the same calibration protocol, now following the normal readout mode with higher frequency. The selected reference f0-g1 reset for our system is highlighted in red and uses $\Omega_{f0g1}/2\pi = 350$ MHz and the higher normal mode of the coupled resonator-filter subsystem.

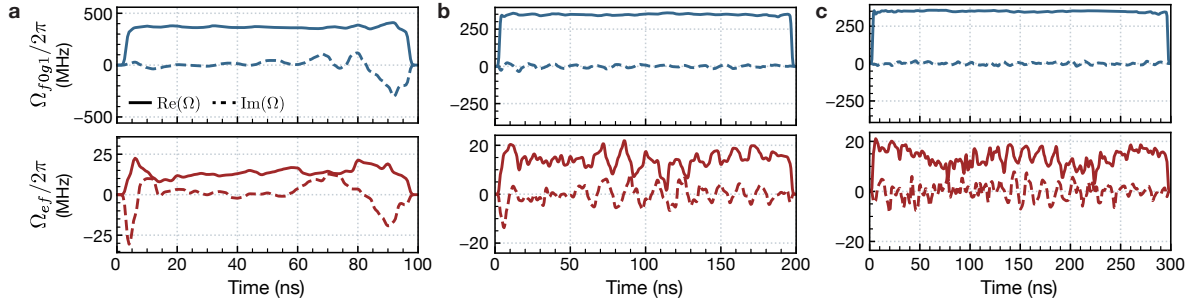


Figure S4. Optimal pulse shapes for reset durations of (a) 100 ns, (b) 200 ns, and (c) 300 ns. The f0-g1 drive Ω_{f0g1} is illustrated in blue and the e-f drive Ω_{ef} in red. Recall that although the features of Ω_{ef} appear to be sharp, the pulses are filtered using an experimentally conservative gaussian filter with a bandwidth of 250 MHz, here and in all of our simulations. This realistic filtering makes these pulses easily implementable in practice.

6.4 Perspectives

La méthodologie de contrôle optimal quantique dans de grands systèmes quantiques ouverts développée dans ce chapitre, de même que les résultats de mesure et d’initialisation d’un transmon, ouvrent la porte à plusieurs avenues de recherche. Tout d’abord, il serait très intéressant de démontrer expérimentalement l’utilisation des contrôles optimisés afin de réaliser la mesure et l’initialisation d’un qubit supraconducteur. Le gain en performance associé à ces opérations optimisées pourrait s’avérer très important vers la réalisation d’un ordinateur quantique tolérant aux fautes. Du côté des explorations numériques d’optimisation de processus quantiques, je discute dans ce qui suit de quelques applications à court et moyen terme que je trouve particulièrement prometteuses. Avant cela, je présente brièvement le programme `DYNAMIQS`, un outil permettant de grandement faciliter ce genre de projets.

6.4.1 Librairie `DYNAMIQS`

Réalisant la grande utilité et applicabilité de notre approche de simulation de grands systèmes quantiques ouverts et du calcul de gradients arbitraires sur ces dynamiques, ainsi que le besoin dans la communauté des circuits supraconducteurs, nous avons décidé de créer une librairie en langage Python partagée librement avec un code source ouvert. Avec pour objectif d’accélérer la simulation et l’optimisation de dynamiques quantiques, nous avons nommé cette librairie `DYNAMIQS` [79]. Bien que toujours en développement, `DYNAMIQS` possède déjà une communauté active avec plus d’une centaine d’utilisateurs et de nombreux contributeurs. L’ensemble du code est aujourd’hui majoritairement développé par Pierre Guilmin, Ronan Gautier et Adrien Bocquet. Ma principale contribution se trouve dans l’élaboration initiale de la librairie ainsi que dans les nombreuses discussions sur les caractéristiques techniques du code.

Développé en Python, `DYNAMIQS` offre une interface similaire à la très populaire librairie QuTiP [68], mais présente quatre avantages importants provenant principalement de l’utilisation de techniques développées en apprentissage automatique. Premièrement, en utilisant la méthode de l’adjoint ainsi que ses variations numériques [197], il est possible d’obtenir le gradient des dynamiques simulées, et ce, pour des fonctions de coût et des paramètres arbitraires. L’utilisation de l’approche développée dans l’article de ce chapitre permet de rester très efficace au niveau de la mémoire numérique nécessaire et ainsi d’être en mesure de simuler et d’optimiser la dynamique de grands systèmes quantiques ouverts.

Deuxièmement, `DYNAMIQS` offre d'énormes accélérations pour la simulation d'un système où l'on s'intéresse aux dynamiques utilisant de nombreux paramètres différents. En effet, tous les paramètres de simulation tels que l'Hamiltonien du système et les opérateurs de dissipations peuvent être mis en lot et résolus simultanément. Cette approche est communément appelée *batching* en apprentissage automatique et permet essentiellement de réaliser de façon vectorielle ce qui est normalement exécuté avec une boucle séquentielle. Troisièmement, l'ensemble des simulations et des opérations numériques peuvent être accélérées à l'aide de processeurs graphiques (GPUs) sans effort additionnel. Ces processeurs excellent entre autres à la multiplication de grandes matrices, ce qui correspond exactement aux opérations utilisées lors de la résolution de l'équation maîtresse. Finalement, alors que le code Python est reconnu pour être significativement plus lent que d'autres langages, `DYNAMIQS` est basé sur la librairie JAX, laquelle réalise une compilation à la volée du code Python vers un code optimisé au niveau de la machine [78]. Cette compilation rend nos algorithmes très rapides malgré leur formulation accessible en Python.

Somme toute, `DYNAMIQS` ouvre la porte à de multiples applications en contrôle optimal et en caractérisation de systèmes quantiques (ou estimation de paramètres), en plus d'offrir diverses accélérations dans la simulation de dynamiques quantiques. Des extraits de code présentant l'interface de `DYNAMIQS` et son utilisation dans la simulation et l'optimisation des dynamiques d'un oscillateur harmonique sont présentés à la Fig. 6.1.

6.4.2 Exemple d'une application directe de la méthode développée

Parmi les multiples applications possibles de la méthodologie développée ici, par exemple en utilisant `DYNAMIQS`, je considère que l'une des plus prometteuses concerne l'optimisation des opérations nécessaires à la correction d'erreur quantique dans des systèmes bosoniques [208]. Dans les réalisations en circuits supraconducteurs de tels codes, l'information quantique est typiquement contenue dans les degrés de liberté d'une cavité micro-onde, laquelle est couplée à un système non linéaire tel qu'un transmon ou un circuit comprenant des jonctions Josephson. Afin de préparer, préserver, manipuler et mesurer cette information quantique encodée, il existe de nombreux protocoles appliquant des impulsions micro-ondes à la cavité ainsi qu'à son système auxiliaire non linéaire. Alors que ces contrôles sont généralement très simples et de longues durées, l'utilisation d'une approche de contrôle optimal permettrait potentiellement d'utiliser davantage la structure de ces systèmes afin de diminuer l'erreur logique résultant de ces opérations.

Dans une réalisation du code de Gottesman-Kitaev-Preskill (GKP) [209] par exemple, l'information sur l'état logique du système est extraite à l'aide d'un protocole spécialisé

```

import dynamiqs as dq
import jax.numpy as jnp

# Paramètres
n = 128 # Dimension de l'espace d'Hilbert
omega = 1.0 # Fréquence
kappa = 0.1 # Taux de relaxation
alpha = 1.0 # Amplitude de l'état cohérent initial

# Hamiltonien et opérateurs de dissipation
a = dq.destroy(n)
H = omega * dq.dag(a) @ a
jump_ops = [jnp.sqrt(kappa) * a]
# Champ d'application de la simulation
psi0 = dq.coherent(n, alpha)
tsave = jnp.linspace(0, 1.0, 101)

# Exécution de la simulation
result = dq.mesolve(H, jump_ops, psi0, tsave)
print(result)

|██████████| 100.0% ♦ elapsed 66.94ms ♦ remaining 0.00ms
==== MESolveResult ====
Solver : Tsit5
Infos : 7 steps (7 accepted, 0 rejected)
States : Array complex64 (101, 128, 128) | 12.62 Mb

```

(a)

```

import dynamiqs as dq
import jax.numpy as jnp
import jax

# Paramètres
n = 128 # Dimension de l'espace d'Hilbert
omega = 1.0 # Fréquence
kappa = 0.1 # Taux de relaxation
alpha = 1.0 # Amplitude de l'état cohérent initial

def population(omega, kappa, alpha):
    """Retourne la population finale de l'oscillateur."""
    # Hamiltonien et opérateurs de dissipation
    a = dq.destroy(n)
    H = omega * dq.dag(a) @ a
    jump_ops = [jnp.sqrt(kappa) * a]
    # Champ d'application de la simulation
    psi0 = dq.coherent(n, alpha)
    tsave = jnp.linspace(0, 1.0, 101)

    # Exécution de la simulation
    result = dq.mesolve(H, jump_ops, psi0, tsave)

    return dq.expect(dq.number(n), result.states[-1]).real

# Calcul des gradients par rapport à omega, kappa et alpha
grad_population = jax.grad(population, argnums=(0, 1, 2))
grads = grad_population(omega, kappa, alpha)
print(f'Gradient p.r. à omega={grads[0]:.2f}')
print(f'Gradient p.r. à kappa={grads[1]:.2f}')
print(f'Gradient p.r. à alpha={grads[2]:.2f}')

```

(c)

```

|██████████| 100.0% ♦ elapsed 86.63ms ♦ remaining 0.00ms
Gradient p.r. à omega=0.00
Gradient p.r. à kappa=-0.90
Gradient p.r. à alpha=1.81

```

(b)

(d)

FIGURE 6.1 – Extraits de code Python utilisant la librairie `DYNAMIQS` disponible en libre accès [79]. (a) Simulation de l'évolution d'un oscillateur harmonique $\hat{H} = \omega \hat{a}^\dagger \hat{a}$ avec une relaxation $\hat{L} = \sqrt{\kappa} \hat{a}$ en résolvant l'équation maîtresse de Lindblad par la fonction `mesolve`. (b) Résultat fourni en sortie permettant de facilement suivre l'évolution de la simulation et d'adapter les options de l'algorithme si nécessaire. (c) Calcul des gradients du nombre de photons final dans la cavité $\bar{n} = \text{Tr}[\hat{a}^\dagger \hat{a} \hat{\rho}(t_f)]$ par rapport à la fréquence de la cavité ω , au taux de relaxation κ et à l'amplitude de l'état cohérent initial α . Cette population pourrait représenter la fonction de coût que l'on souhaite optimiser. Dans ce cas, il suffit de fournir les gradients à un algorithme de descente du gradient, tel qu'Adam [207], afin de mettre à jour les paramètres du système et d'optimiser la métrique voulue. (d) Sortie du code présenté en (c).

afin de stabiliser le système dans son état logique de manière autonome [210]. Ce protocole combine une série d'interactions contrôlées entre la cavité et le système auxiliaire à des réinitialisations répétées de ce système auxiliaire, une technique pouvant être considérée comme de l'ingénierie de la dissipation [211]. Cependant, ce protocole ne réalise pas l'opération optimale pour récupérer l'ensemble de l'information accessible afin de stabiliser l'état GKP [212, 213]. Sachant qu'un gain important est atteignable en principe, il serait alors très intéressant d'utiliser une approche de contrôle optimal quantique en système ouvert afin d'optimiser les contrôles appliqués sur ce système et ainsi d'améliorer la performance du code.

Conclusion

Dans cette thèse, nous avons exploré différentes façons de construire un *bon modèle* pour décrire les dynamiques quantiques de qubits supraconducteurs avec un niveau de précision suffisant pour comprendre les limitations principales des expériences réalisées actuellement. À partir de caractérisations expérimentales détaillées nous permettant de construire ces bons modèles, nous avons également optimisé les façons dont on contrôle ces systèmes complexes afin de réaliser des opérations quantiques avec de meilleures fidélités, notamment des portes logiques et des mesures projectives sur les qubits. Nos méthodologies sont au cœur du développement des technologies quantiques alors qu'elles permettent de caractériser un système quantique, de concevoir un modèle à partir de cette information et finalement d'utiliser ce modèle afin de contrôler l'opération du système avec une erreur minimale. Bien que la plupart des approches développées ici aient été spécialisées à des systèmes de circuits supraconducteurs, elles peuvent aisément être adaptées à de nombreux autres systèmes quantiques d'aujourd'hui et de demain.

Suite à une introduction aux concepts clés de modélisation et de simulation des opérations effectuées avec les qubits supraconducteurs au chapitre 2, nous avons appliqué ces idées à la description détaillée d'une expérience de correction d'erreurs quantiques utilisant 17 qubits au chapitre 3. Nous avons vu spécifiquement les compromis nécessaires entre une modélisation précise des dynamiques physiques et une simulation pouvant être mise à l'échelle à des dispositifs de grande taille. Alors que ces deux types de simulations sont essentiels au développement des technologies quantiques, des études approfondies sont nécessaires afin de réconcilier ces approches et ainsi prendre en compte les sources de bruit réalistes lors des extrapolations à grande échelle.

Au chapitre 4, nous avons exploré une méthodologie alternative de caractérisation et de contrôle basée sur l'apprentissage automatique. Nous avons démontré l'utilité d'une telle approche sur le contrôle d'un qubit supraconducteur et décrit de nombreuses applications de cette méthode qui promettent d'être utiles pour améliorer la fidélité et la rapidité des

opérations réalisées avec nos dispositifs quantiques. Plus généralement, le codéveloppement de l'apprentissage automatique et des technologies quantiques ne fait que commencer et les efforts dans cette direction présentent un potentiel énorme.

Au chapitre 5, nous avons introduit deux protocoles de caractérisation quantique, CAFE et MEADD, qui sont spécifiquement conçus afin d'extraire l'information utile au contrôle des opérations quantiques de manière précise et efficace. Il s'avère que pour atteindre des fidélités de portes logiques supérieures à 99.9 % pour des centaines de qubits différents, il est nécessaire d'utiliser des protocoles robustes qui sont spécialisés au système et même aux séquences d'opérations étudiés. Avec la quête vers des fidélités et des systèmes toujours plus grands, je m'attends à ce que de nombreux autres protocoles de ce genre soient développés et largement utilisés.

Au chapitre 6, nous avons finalement développé une méthode numérique d'optimisation des dynamiques quantiques dans de grands systèmes ouverts et avons démontré son utilité dans la conception des contrôles utilisés dans la mesure dispersive et la réinitialisation du transmon. Ces résultats démontrent que des gains d'envergure en fidélité de contrôle sont possibles lorsque l'on possède un bon modèle des dynamiques du dispositif. Plus généralement, l'ensemble des travaux de cette thèse mène à cette même conclusion selon laquelle la construction d'un bon modèle est essentielle à l'amélioration des technologies quantiques. Alors que nous nous sommes concentrés sur les portes logiques à un et deux qubits et sur la mesure et la réinitialisation d'un transmon, les outils de caractérisation et de contrôle que nous avons développés peuvent être appliqués à de nombreuses approches prometteuses en détection quantique [214], en ingénierie de la dissipation [185], en correction d'erreurs quantiques bosonique [208], en communication quantique [215] et bien plus.

Prenant un dernier pas de recul sur l'ensemble des travaux présentés dans cette thèse, je considère qu'étant donné la complexité des systèmes que l'on cherche à contrôler ainsi que le besoin de constamment réduire les erreurs des opérations quantiques réalisées, le travail de modélisation, de caractérisation et de contrôle sera toujours à parfaire à mesure que les systèmes évoluent. De la même façon que mes travaux s'appuient sur les résultats de nombreuses années de recherche réalisées par la communauté, j'ai espoir que mes contributions au domaine pourront être utiles au développement des technologies quantiques de demain, et je peux me reconforter de les voir être utilisées aujourd'hui même [216, 217, 218, 15]. La recherche dans le domaine de l'informatique quantique s'avère toujours pleine de surprises fascinantes nous confrontant continuellement à notre ignorance, mais nous récompensant toujours pour nos efforts de compréhension d'une manière concrète. Je me sens très privilégié d'avoir la chance de prendre part à cette aventure scientifique et je suis enthousiaste de découvrir ce que son futur nous réserve.

Bibliographie

- [1] S. KRINNER, N. LACROIX, A. REMM, A. DI PAOLO, E. GENOIS, C. LEROUX, C. HEL-
LINGS, S. LAZAR, F. SWIADEK, J. HERRMANN, G. J. NORRIS, C. K. ANDERSEN,
M. MÜLLER, A. BLAIS, C. EICHLER et A. WALLRAFF, « Realizing repeated quantum
error correction in a distance-three surface code », *Nature*, vol. 605, p. 669–674, mai
2022. (Cité aux pages [vii](#), [21](#), [26](#), [36](#), [37](#), [94](#), and [97](#).)
- [2] É. GENOIS, N. J. STEVENSON et A. BLAIS, « Quantum optimal control of supercon-
ducting qubits based on machine-learning characterization ». En préparation, 2024.
(Cité aux pages [viii](#), [53](#), and [65](#).)
- [3] D. M. DEBROY, É. GENOIS, J. A. GROSS, W. MRUCZKIEWICZ, K. LEE, S. HONG,
Z. CHEN, V. SMELYANSKIY et Z. JIANG, « Context-aware fidelity estimation », *Phys.*
Rev. Res., vol. 5, p. 043202, Dec 2023. (Cité aux pages [viii](#), [94](#), and [104](#).)
- [4] J. A. GROSS, É. GENOIS, D. M. DEBROY, Y. ZHANG, W. MRUCZKIEWICZ, Z.-P. CIAN
et Z. JIANG, « Characterizing coherent errors using matrix-element amplification »,
2024. arXiv :2404.12550. (Cité aux pages [viii](#), [88](#), [100](#), and [122](#).)
- [5] R. GAUTIER, É. GENOIS et A. BLAIS, « Optimal control in large open quantum systems :
the case of transmon readout and reset », 2024. arXiv :2403.14765. (Cité aux pages [viii](#)
and [151](#).)
- [6] S. J. GLASER, U. BOSCAIN, T. CALARCO, C. P. KOCH, W. KÖCKENBERGER, R. KOSLOFF,
I. KUPROV, B. LUY, S. SCHIRMER, T. SCHULTE-HERBRÜGGEN, D. SUGNY et F. K.
WILHELM, « Training schrödinger’s cat : quantum optimal control : Strategic report
on current status, visions and goals for research in europe », *The European Physical*
Journal D, vol. 69, déc. 2015. (Cité aux pages [1](#) and [54](#).)
- [7] « 40 years of quantum computing », *Nature Reviews Physics*, vol. 4, p. 1–1, jan. 2022.
Publisher : Nature Publishing Group. (Cité à la page [1](#).)
- [8] T. D. LADD, F. JELEZKO, R. LAFLAMME, Y. NAKAMURA, C. MONROE et J. L. O’BRIEN,
« Quantum computers », *Nature*, vol. 464, p. 45–53, mars 2010. Publisher : Nature
Publishing Group. (Cité à la page [1](#).)
- [9] J. PRESKILL, « Quantum computing in the nist era and beyond », *Quantum*, vol. 2,
p. 79, août 2018. (Cité aux pages [1](#) and [22](#).)
- [10] P. W. SHOR, « Scheme for reducing decoherence in quantum computer memory »,
Phys. Rev. A, vol. 52, p. R2493–R2496, Oct 1995. (Cité à la page [2](#).)
- [11] A. KITAEV, « Fault-tolerant quantum computation by anyons », *Annals of Physics*,
vol. 303, p. 2–30, jan. 2003. (Cité à la page [2](#).)

- [12] M. H. DEVORET et R. J. SCHOELKOPF, « Superconducting circuits for quantum information : An outlook », *Science*, vol. 339, no. 6124, p. 1169–1174, 2013. (Cité aux pages 2 and 6.)
- [13] A. BLAIS, A. L. GRIMSMO, S. M. GIRVIN et A. WALLRAFF, « Circuit quantum electrodynamics », *Rev. Mod. Phys.*, vol. 93, p. 025005, May 2021. (Cité aux pages 2, 6, 7, 9, 10, 11, 12, 14, 15, 17, 19, 89, 97, and 149.)
- [14] Y. KIM, A. EDDINS, S. ANAND, K. X. WEI, E. V. D. BERG, S. ROSENBLATT, H. NAYFEH, Y. WU, M. ZALETEL, K. TEMME et A. KANDALA, « Evidence for the utility of quantum computing before fault tolerance », *Nature*, vol. 618, p. 500–505, 2023. (Cité aux pages 2 and 10.)
- [15] R. ACHARYA, L. AGHABABAIE-BENI, I. ALEINER, T. I. ANDERSEN, M. ANSMANN, F. ARUTE, K. ARYA, A. ASFAW, N. ASTRAKHANTSEV, J. ATALAYA, R. BABBUSH, D. BACON, B. BALLARD, J. C. BARDIN, J. BAUSCH, A. BENGTSSON, A. BILMES, S. BLACKWELL, S. BOIXO *et al.*, « Quantum error correction below the surface code threshold », 2024. arXiv :2408.13687. (Cité aux pages 2, 10, 94, 97, and 172.)
- [16] A. G. FOWLER, M. MARIANTONI, J. M. MARTINIS et A. N. CLELAND, « Surface codes : Towards practical large-scale quantum computation », *Physical Review A*, vol. 86, p. 032324, sept. 2012. Publisher : American Physical Society. (Cité aux pages 2, 22, and 23.)
- [17] E. T. CAMPBELL, B. M. TERHAL et C. VUILLOT, « Roads towards fault-tolerant universal quantum computation », *Nature*, vol. 549, p. 172–179, sept. 2017. Publisher : Nature Publishing Group. (Cité à la page 2.)
- [18] G. E. P. BOX, « Science and statistics », *Journal of the American Statistical Association*, vol. 71, no. 356, p. 791–799, 1976. (Cité à la page 2.)
- [19] J. BARDEEN, L. N. COOPER et J. R. SCHRIEFFER, « Theory of superconductivity », *Phys. Rev.*, vol. 108, p. 1175–1204, Dec 1957. (Cité à la page 5.)
- [20] U. ECKERN, G. SCHÖN et V. AMBEGAOKAR, « Quantum dynamics of a superconducting tunnel junction », *Phys. Rev. B*, vol. 30, p. 6419–6431, Dec 1984. (Cité à la page 5.)
- [21] M. DEVORET, B. HUARD, R. SCHOELKOPF et L. F. CUGLIANDOLO, *Quantum Machines : Measurement and Control of Engineered Quantum Systems : Lecture Notes of the Les Houches Summer School : Volume 96, July 2011*. Oxford University Press, 06 2014. (Cité à la page 5.)
- [22] P. KRANTZ, M. KJAERGAARD, F. YAN, T. P. ORLANDO, S. GUSTAVSSON et W. D. OLIVER, « A Quantum Engineer’s Guide to Superconducting Qubits », *Applied Physics Reviews*, vol. 6, p. 021318, juin 2019. (Cité aux pages 6, 13, 17, and 51.)
- [23] U. VOOL et M. DEVORET, « Introduction to quantum electromagnetic circuits », *International Journal of Circuit Theory and Applications*, vol. 45, p. 897–934, juin 2017. (Cité à la page 6.)
- [24] M. KJAERGAARD, M. E. SCHWARTZ, J. BRAUMÜLLER, P. KRANTZ, J. I.-J. WANG, S. GUSTAVSSON et W. D. OLIVER, « Superconducting qubits : Current state of play », *Annual Review of Condensed Matter Physics*, vol. 11, p. 369–395, mars 2020. (Cité à la page 6.)

- [25] J. CLARKE et F. K. WILHELM, « Superconducting quantum bits », *Nature*, vol. 453, p. 1031–1042, 2008. (Cité à la page 6.)
- [26] C. R. H. McRAE, H. WANG, J. GAO, M. R. VISSERS, T. BRECHT, A. DUNSWORTH, D. P. PAPPAS et J. MUTUS, « Materials loss measurements using superconducting microwave resonators », *Review of Scientific Instruments*, vol. 91, sept. 2020. (Cité aux pages 6 and 7.)
- [27] J. KOCH, T. M. YU, J. GAMBETTA, A. A. HOUCK, D. I. SCHUSTER, J. MAJER, A. BLAIS, M. H. DEVORET, S. M. GIRVIN et R. J. SCHOELKOPF, « Charge-insensitive qubit design derived from the cooper pair box », *Phys. Rev. A*, vol. 76, p. 042319, Oct 2007. (Cité aux pages 6, 9, 10, 12, 14, 27, and 148.)
- [28] A. BLAIS, R.-S. HUANG, A. WALLRAFF, S. M. GIRVIN et R. J. SCHOELKOPF, « Cavity quantum electrodynamics for superconducting electrical circuits : an architecture for quantum computation », *Physical Review A*, vol. 69, p. 062320, juin 2004. (Cité aux pages 6, 11, and 12.)
- [29] D. M. POZAR, *Microwave engineering; 3rd ed.* Hoboken, NJ : Wiley, 2005. (Cité à la page 7.)
- [30] H. CARMICHAEL, *Statistical Methods in Quantum Optics 1 : Master Equations and Fokker-Planck Equations*. Physics and astronomy online library, Springer, 1998. (Cité aux pages 8 and 19.)
- [31] B. JOSEPHSON, « Possible new effects in superconductive tunnelling », *Physics Letters*, vol. 1, no. 7, p. 251–253, 1962. (Cité à la page 8.)
- [32] M. H. DEVORET, « Quantum fluctuations in electrical circuits », in *Quantum Fluctuations, Les Houches, Session LXIII* (S. REYNAUD, E. GIACOBINO et J. ZINN-JUSTIN, édés), chap. 10, Elsevier Science, 1995. (Cité à la page 8.)
- [33] K. SERNIAK, *Nonequilibrium Quasiparticles in Superconducting Qubits*. Thèse doctorat, Yale University, 2019. (Cité à la page 9.)
- [34] J. CLARKE, A. N. CLELAND, M. H. DEVORET, D. ESTEVE et J. M. MARTINIS, « Quantum mechanics of a macroscopic variable : The phase difference of a josephson junction », *Science*, vol. 239, no. 4843, p. 992–997, 1988. (Cité à la page 10.)
- [35] A. WALLRAFF, D. I. SCHUSTER, A. BLAIS, L. FRUNZIO, R.-S. HUANG, J. MAJER, S. KUMAR, S. M. GIRVIN et R. J. SCHOELKOPF, « Circuit Quantum Electrodynamics : Coherent Coupling of a Single Photon to a Cooper Pair Box », *Nature*, vol. 431, p. 162–167, sept. 2004. arXiv : cond-mat/0407325. (Cité à la page 11.)
- [36] H. WISEMAN et G. MILBURN, *Quantum Measurement and Control*. Cambridge University Press, 2010. (Cité aux pages 13, 19, 54, and 89.)
- [37] D. I. SCHUSTER, A. WALLRAFF, A. BLAIS, L. FRUNZIO, R.-S. HUANG, J. MAJER, S. M. GIRVIN et R. J. SCHOELKOPF, « ac stark shift and dephasing of a superconducting qubit strongly coupled to a cavity field », *Phys. Rev. Lett.*, vol. 94, p. 123602, Mar 2005. (Cité à la page 13.)
- [38] A. BARENCO, C. H. BENNETT, R. CLEVE, D. P. DiVINCENZO, N. MARGOLUS, P. SHOR, T. SLEATOR, J. A. SMOLIN et H. WEINFURTER, « Elementary gates for quantum computation », *Physical Review A*, vol. 52, p. 3457–3467, nov. 1995. (Cité à la page 14.)

- [39] R. L. FAGALY, « Superconducting quantum interference device instruments and applications », *Review of Scientific Instruments*, vol. 77, p. 101101, 10 2006. (Cité à la page 14.)
- [40] J. M. CHOW, L. DICARLO, J. M. GAMBETTA, F. MOTZOI, L. FRUNZIO, S. M. GIRVIN et R. J. SCHOELKOPF, « Optimized driving of superconducting artificial atoms for improved single-qubit gates », *Phys. Rev. A*, vol. 82, p. 040305, Oct 2010. (Cité à la page 17.)
- [41] F. MOTZOI, J. M. GAMBETTA, P. REBENTROST et F. K. WILHELM, « Simple Pulses for Elimination of Leakage in Weakly Nonlinear Qubits », *Physical Review Letters*, vol. 103, p. 110501, sept. 2009. (Cité aux pages 17 and 150.)
- [42] J. M. GAMBETTA, F. MOTZOI, S. T. MERKEL et F. K. WILHELM, « Analytic control methods for high-fidelity unitary operations in a weakly nonlinear oscillator », *Phys. Rev. A*, vol. 83, p. 012308, Jan 2011. (Cité à la page 17.)
- [43] M. A. NIELSEN, « A simple formula for the average gate fidelity of a quantum dynamical operation », *Physics Letters A*, vol. 303, p. 249–252, oct. 2002. (Cité à la page 17.)
- [44] A. HASHIM, L. B. NGUYEN, N. GOSS, B. MARINELLI, R. K. NAIK, T. CHISTOLINI, J. HINES, J. P. MARCEAUX, Y. KIM, P. GOKHALE, T. TOMESH, S. CHEN, L. JIANG, S. FERRACIN, K. RUDINGER, T. PROCTOR, K. C. YOUNG, R. BLUME-KOHOUT et I. SIDDIQI, « A practical introduction to benchmarking and characterization of quantum computers », 2024. arXiv :2408.12064. (Cité aux pages 17 and 96.)
- [45] C. DANKERT, R. CLEVE, J. EMERSON et E. LIVINE, « Exact and approximate unitary 2-designs and their application to fidelity estimation », *Physical Review A*, vol. 80, juil. 2009. (Cité aux pages 17 and 95.)
- [46] D. F. GRIFFITHS et D. J. HIGHAM, *Numerical Methods for Ordinary Differential Equations*. Springer London, nov. 2010. (Cité aux pages 18, 86, and 147.)
- [47] C. GARDINER et P. ZOLLER, *Quantum Noise : A Handbook of Markovian and Non-Markovian Quantum Stochastic Methods with Applications to Quantum Optics*. Springer Series in Synergetics, Springer, 2004. (Cité à la page 19.)
- [48] H.-P. BREUER et F. PETRUCCIONE, *The Theory of Open Quantum Systems*. Oxford University Press, 01 2007. (Cité aux pages 19 and 143.)
- [49] C. RYAN-ANDERSON, J. G. BOHNET, K. LEE, D. GRESH, A. HANKIN, J. P. GAEBLER, D. FRANCOIS, A. CHERNOGUZOV, D. LUCCHETTI, N. C. BROWN, T. M. GATTERMAN, S. K. HALIT, K. GILMORE, J. A. GERBER, B. NEYENHUIS, D. HAYES et R. P. STUTZ, « Realization of real-time fault-tolerant quantum error correction », *Phys. Rev. X*, vol. 11, p. 041058, Dec 2021. (Cité aux pages 21 and 83.)
- [50] M. REIHER, N. WIEBE, K. M. SVORE, D. WECKER et M. TROYER, « Elucidating reaction mechanisms on quantum computers », *Proceedings of the National Academy of Sciences*, vol. 114, no. 29, p. 7555–7560, 2017. (Cité à la page 22.)
- [51] C. GIDNEY et M. EKERA, « How to factor 2048 bit rsa integers in 8 hours using 20 million noisy qubits », *Quantum*, vol. 5, p. 433, avril 2021. (Cité à la page 22.)
- [52] R. BABBUSH, J. R. MCCLEAN, M. NEWMAN, C. GIDNEY, S. BOIXO et H. NEVEN, « Focus beyond quadratic speedups for error-corrected quantum advantage », *PRX Quantum*, vol. 2, p. 010103, Mar 2021. (Cité à la page 22.)

- [53] É. GOUZIEN, D. RUIZ, F.-M. LE RÉGENT, J. GUILLAUD et N. SANGOUARD, « Performance analysis of a repetition cat code architecture : Computing 256-bit elliptic curve logarithm in 9 hours with 126133 cat qubits », *Physical Review Letters*, vol. 131, juil. 2023. (Cité à la page 22.)
- [54] D. GOTTESMAN, « Stabilizer codes and quantum error correction », 1997. arXiv :9705052. (Cité à la page 22.)
- [55] V. V. SIVAK, A. EICKBUSCH, B. ROYER, S. SINGH, I. TSIOUTSIOS, S. GANJAM, A. MIANO, B. L. BROCK, A. Z. DING, L. FRUNZIO, S. M. GIRVIN, R. J. SCHOELKOPF et M. H. DEVORET, « Real-time quantum error correction beyond break-even », *Nature*, vol. 616, p. 50–55, mars 2023. (Cité aux pages 22, 60, and 62.)
- [56] U. RÉGLADE, A. BOCQUET, R. GAUTIER, J. COHEN, A. MARQUET, E. ALBERTINALE, N. PANKRATOVA, M. HALLÉN, F. RAUTSCHKE, L.-A. SELLEM, P. ROUCHON, A. SARLETTE, M. MIRRAHIMI, P. CAMPAGNE-IBARCQ, R. LESCOANNE, S. JEZOUIN et Z. LEGHTAS, « Quantum control of a cat qubit with bit-flip times exceeding ten seconds », *Nature*, vol. 629, p. 778–783, mai 2024. (Cité à la page 22.)
- [57] E. DENNIS, A. KITAEV, A. LANDAHL et J. PRESKILL, « Topological quantum memory », *Journal of Mathematical Physics*, vol. 43, p. 4452–4505, sept. 2002. (Cité à la page 23.)
- [58] S. B. BRAVYI et A. Y. KITAEV, « Quantum codes on a lattice with boundary », 1998. arXiv :9811052. (Cité à la page 23.)
- [59] S. AARONSON et D. GOTTESMAN, « Improved simulation of stabilizer circuits », *Physical Review A*, vol. 70, nov. 2004. (Cité à la page 25.)
- [60] C. GIDNEY, « Stim : a fast stabilizer circuit simulator », *Quantum*, vol. 5, p. 497, juil. 2021. (Cité à la page 25.)
- [61] GOOGLE QUANTUM AI, « qsim », sept. 2020. <https://quantumai.google/qsim>. (Cité aux pages 25 and 26.)
- [62] M. A. NIELSEN et I. L. CHUANG, *Quantum Computation and Quantum Information : 10th Anniversary Edition*. Cambridge University Press, 2011. (Cité aux pages 25 and 96.)
- [63] GOOGLE QUANTUM AI, « Suppressing quantum errors by scaling a surface code logical qubit », *Nature*, vol. 614, 2023. (Cité aux pages 25, 26, and 37.)
- [64] S. KRINNER, S. LAZAR, A. REMM, C. ANDERSEN, N. LACROIX, G. NORRIS, C. HELINGS, M. GABUREAC, C. EICHLER et A. WALLRAFF, « Benchmarking coherent errors in controlled-phase gates due to spectator qubits », *Phys. Rev. Appl.*, vol. 14, p. 024042, Aug 2020. (Cité aux pages 26, 33, and 87.)
- [65] J. DALIBARD, Y. CASTIN et K. MØLMER, « Wave-function approach to dissipative processes in quantum optics », *Phys. Rev. Lett.*, vol. 68, p. 580–583, Feb 1992. (Cité à la page 29.)
- [66] K. MØLMER, Y. CASTIN et J. DALIBARD, « Monte carlo wave-function method in quantum optics », *J. Opt. Soc. Am. B*, vol. 10, p. 524–538, Mar 1993. (Cité à la page 29.)
- [67] R. DUM, P. ZOLLER et H. RITSCH, « Monte carlo simulation of the atomic master equation for spontaneous emission », *Phys. Rev. A*, vol. 45, p. 4879–4887, Apr 1992. (Cité à la page 29.)

- [68] J. JOHANSSON, P. NATION et F. NORI, « Qutip 2 : A python framework for the dynamics of open quantum systems », *Computer Physics Communications*, vol. 184, p. 1234–1240, avril 2013. (Cité aux pages 29, 48, and 168.)
- [69] C. K. ANDERSEN, A. REMM, S. LAZAR, S. KRINNER, N. LACROIX, G. J. NORRIS, M. GABUREAC, C. EICHLER et A. WALLRAFF, « Repeated quantum error detection in a surface code », *Nature Physics*, vol. 16, p. 875–880, juin 2020. (Cité aux pages 30 and 31.)
- [70] M. A. ROL, F. BATTISTEL, F. K. MALINOWSKI, C. C. BULTINK, B. M. TARASINSKI, R. VOLLMER, N. HAIDER, N. MUTHUSUBRAMANIAN, A. BRUNO, B. M. TERHAL et L. DiCARLO, « Fast, high-fidelity conditional-phase gate exploiting leakage interference in weakly anharmonic superconducting qubits », *Phys. Rev. Lett.*, vol. 123, p. 120502, Sep 2019. (Cité aux pages 32, 52, and 94.)
- [71] V. NEGÎRNEAC, H. ALI, N. MUTHUSUBRAMANIAN, F. BATTISTEL, R. SAGASTIZABAL, M. S. MOREIRA, J. F. MARQUES, W. J. VLOTHUIZEN, M. BEEKMAN, C. ZACHARIADIS, N. HAIDER, A. BRUNO et L. DiCARLO, « High-fidelity controlled-z gate with maximal intermediate leakage operating at the speed limit in a superconducting quantum processor », *Phys. Rev. Lett.*, vol. 126, p. 220502, Jun 2021. (Cité aux pages 32, 52, and 88.)
- [72] T. E. O'BRIEN, B. TARASINSKI et L. DiCARLO, « Density-matrix simulation of small surface codes under current and projected experimental noise », *npj Quantum Information*, vol. 3, sept. 2017. (Cité à la page 35.)
- [73] C. HUANG, X. NI, F. ZHANG, M. NEWMAN, D. DING, X. GAO, T. WANG, H.-H. ZHAO, F. WU, G. ZHANG, C. DENG, H.-S. KU, J. CHEN et Y. SHI, « Alibaba cloud quantum development platform : Surface code simulations with crosstalk », 2020. arXiv :2002.08918. (Cité à la page 35.)
- [74] J. MARSHALL et D. KAFRI, « Incoherent approximation of leakage in quantum error correction », 2023. arXiv :2312.10277. (Cité à la page 47.)
- [75] B. M. VARBANOV, F. BATTISTEL, B. M. TARASINSKI, V. P. OSTROUKH, T. E. O'BRIEN, L. DiCARLO et B. M. TERHAL, « Leakage detection for a transmon-based surface code », *npj Quantum Information*, vol. 6, déc. 2020. (Cité à la page 47.)
- [76] J. F. MARQUES, H. ALI, B. M. VARBANOV, M. FINKEL, H. M. VEEN, S. L. M. van der MEER, S. VALLES-SANCLEMENTE, N. MUTHUSUBRAMANIAN, M. BEEKMAN, N. HAIDER, B. M. TERHAL et L. DiCARLO, « All-microwave leakage reduction units for quantum error correction with superconducting transmon qubits », *Phys. Rev. Lett.*, vol. 130, p. 250602, Jun 2023. (Cité à la page 48.)
- [77] N. LACROIX, L. HOFELE, A. REMM, O. BENHAYOUNE-KHADRAOUI, A. McDONALD, R. SHILLITO, S. LAZAR, C. HELLINGS, F. SWIADEK, D. COLAO-ZANUZ, A. FLASBY, M. B. PANAH, M. KERSCHBAUM, G. J. NORRIS, A. BLAIS, A. WALLRAFF et S. KRINNER, « Fast flux-activated leakage reduction for superconducting quantum circuits », 2023. arXiv :2309.07060. (Cité à la page 48.)
- [78] J. BRADBURY, R. FROSTIG, P. HAWKINS, M. J. JOHNSON, C. LEARY, D. MACLAURIN, G. NECULA, A. PASZKE, J. VANDERPLAS, S. WANDERMAN-MILNE et Q. ZHANG, « JAX : composable transformations of Python+NumPy programs », 2018. <http://github.com/google/jax>. (Cité aux pages 48, 86, and 169.)

- [79] P. GUILMIN, R. GAUTIER, A. BOCQUET et É. GENOIS, « dynamiqs : an open-source python library for gpu-accelerated and differentiable simulation of quantum systems », dynamiqs.org, 2024. (Cité aux pages 48, 85, 86, 142, 168, and 170.)
- [80] J. M. MARTINIS, K. B. COOPER, R. McDERMOTT, M. STEFFEN, M. ANSMANN, K. D. OSBORN, K. CİCAK, S. OH, D. P. PAPPAS, R. W. SIMMONDS et C. C. YU, « Decoherence in josephson qubits from dielectric loss », *Phys. Rev. Lett.*, vol. 95, p. 210503, Nov 2005. (Cité à la page 50.)
- [81] A. NERSISYAN, S. POLETO, N. ALIDOUST, R. MANENTI, R. RENZAS, C.-V. BUI, K. VU, T. WHYLAND, Y. MOHAN, E. A. SETE, S. STANWYCK, A. BESTWICK et M. REAGOR, « Manufacturing low dissipation superconducting quantum processors », 2019. arXiv :1901.08042. (Cité à la page 50.)
- [82] P. V. KLIMOV, J. KELLY, Z. CHEN, M. NEELEY, A. MEGRANT, B. BURKETT, R. BARENDTS, K. ARYA, B. CHIARO, Y. CHEN, A. DUNSWORTH, A. FOWLER, B. FOXEN, C. GIDNEY, M. GIUSTINA, R. GRAFF, T. HUANG, E. JEFFREY, E. LUCERO, J. Y. MUTUS, O. NAAMAN, C. NEILL, C. QUINTANA, P. ROUSHAN, D. SANK, A. VAINSENCHER, J. WENNER, T. C. WHITE, S. BOIXO, R. BABBUSH, V. N. SMELYANSKIY, H. NEVEN et J. M. MARTINIS, « Fluctuations of energy-relaxation times in superconducting qubits », *Phys. Rev. Lett.*, vol. 121, p. 090502, Aug 2018. (Cité aux pages 50 and 93.)
- [83] M. CARROLL, S. ROSENBLATT, P. JURCEVIC, I. LAUER et A. KANDALA, « Dynamics of superconducting qubit relaxation times », *Nature Communications*, vol. 8, no. 132, 2022. (Cité à la page 50.)
- [84] I. SIDDIQI, « Engineering high-coherence superconducting qubits », *Nature Reviews Materials*, vol. 6, p. 875 – 891, 2021. (Cité à la page 50.)
- [85] A. P. M. PLACE, L. V. H. RODGERS, P. MUNDADA, B. M. SMITHAM, M. FITZPATRICK, Z. LENG, A. PREMKUMAR, J. BRYON, A. VRAJITOAREA, S. SUSSMAN, G. CHENG, T. MADHAVAN, H. K. BABLA, X. H. LE, Y. GANG, B. JÄCK, A. GYENIS, N. YAO, R. J. CAVA, N. P. de LEON et A. A. HOUCK, « New material platform for superconducting transmon qubits with coherence times exceeding 0.3 milliseconds », *Nature Communications*, vol. 12, mars 2021. (Cité à la page 50.)
- [86] M. BAL, A. A. MURTHY, S. ZHU, F. CRISA, X. YOU, Z. HUANG, T. ROY, J. LEE, D. V. ZANTEN, R. PILIPENKO, I. NEKRASHEVICH, A. LUNIN, D. BAFIA, Y. KRASNIKOVA et al., « Systematic improvements in transmon qubit coherence enabled by niobium surface encapsulation », *npj Quantum Information*, vol. 10, avril 2024. (Cité à la page 50.)
- [87] C. J. KOPAS, D. P. GORONZY, T. PHAM, C. G. T. CASTANEDO, M. CHENG, R. COCHRANE, P. NAST, E. LACHMAN, N. Z. ZHELEV, A. VALLIERES, A. A. MURTHY, J. SU OH, L. ZHOU, M. J. KRAMER, H. CANSIZOGLU, M. J. BEDZYK, V. P. DRAVID, A. ROMANENKO, A. GRASSELLINO, J. Y. MUTUS, M. C. HERSAM et K. YADAVALLI, « Enhanced superconducting qubit performance through ammonium fluoride etch », 2024. arXiv :2408.02863. (Cité à la page 50.)
- [88] P. V. KLIMOV, J. KELLY, J. M. MARTINIS et H. NEVEN, « The snake optimizer for learning quantum processor control parameters », 2020. arXiv :2006.04594. (Cité aux pages 50 and 89.)
- [89] P. V. KLIMOV, A. BENGTTSSON, C. QUINTANA, A. BOURASSA, S. HONG, A. DUNSWORTH, K. J. SATZINGER, W. P. LIVINGSTON, V. SIVAK, M. Y. NIU, T. I. ANDERSEN,

- Y. ZHANG, D. CHIK, Z. CHEN, C. NEILL, C. ERICKSON, A. GRAJALES DAU, A. MEGRANT, P. ROUSHAN, A. N. KOROTKOV, J. KELLY, V. SMELYANSKIY, Y. CHEN et H. NEVEN, « Optimizing quantum gates towards the scale of logical qubits », *Nature Communications*, vol. 15, mars 2024. (Cit      la page 50.)
- [90] B. SARABI, A. N. RAMANAYAKA, A. L. BURIN, F. C. WELLSTOOD et K. D. OSBORN, « Projected dipole moments of individual two-level defects extracted using circuit quantum electrodynamics », *Phys. Rev. Lett.*, vol. 116, p. 167002, Apr 2016. (Cit      la page 50.)
- [91] G. J. GRABOVSKIJ, T. PEICHL, J. LISENFELD, G. WEISS et A. V. USTINOV, « Strain tuning of individual atomic tunneling systems detected by a superconducting qubit », *Science*, vol. 338, no. 6104, p. 232–234, 2012. (Cit      la page 50.)
- [92] L. V. ABDURAKHIMOV, I. MAHBOOB, H. TOIDA, K. KAKUYANAGI, Y. MATSUZAKI et S. SAITO, « Driven-state relaxation of a coupled qubit-defect system in spin-locking measurements », *Phys. Rev. B*, vol. 102, p. 100502, Sep 2020. (Cit      la page 50.)
- [93] Y. KIM, L. C. G. GOVIA, A. DANE, E. van den BERG, D. M. ZAJAC, B. MITCHELL, Y. LIU, K. BALAKRISHNAN, G. KEEFE, A. STABILE, E. PRITCHETT, J. STEHLIK et A. KANDALA, « Error mitigation with stabilized noise in superconducting quantum processors », 2024. arXiv :2407.02467. (Cit      la page 50.)
- [94] Y. BAUM, M. AMICO, S. HOWELL, M. HUSH, M. LIUZZI, P. MUNDADA, T. MERKH, A. R. CARVALHO et M. J. BIERCUK, « Experimental deep reinforcement learning for error-robust gate-set design on a superconducting quantum computer », *PRX Quantum*, vol. 2, p. 040324, Nov 2021. (Cit   aux pages 51, 60, and 61.)
- [95] S. GUSTAVSSON, O. ZWIER, J. BYLANDER, F. YAN, F. YOSHIHARA, Y. NAKAMURA, T. P. ORLANDO et W. D. OLIVER, « Improving quantum gate fidelities by using a qubit to measure microwave pulse distortions », *Phys. Rev. Lett.*, vol. 110, p. 040502, Jan 2013. (Cit      la page 51.)
- [96] M. JERGER, A. KULIKOV, Z. VASSELIN et A. FEDOROV, « In situ characterization of qubit control lines : A qubit as a vector network analyzer », *Phys. Rev. Lett.*, vol. 123, p. 150501, Oct 2019. (Cit      la page 51.)
- [97] M. A. ROL, L. CIORCIARO, F. K. MALINOWSKI, B. M. TARASINSKI, R. E. SAGASTIZABAL, C. C. BULTINK, Y. SALATHE, N. HAANDBAEK, J. SEDIVY et L. DICARLO, « Time-domain characterization and correction of on-chip distortion of control pulses in a quantum processor », *Applied Physics Letters*, vol. 116, p. 054001, 02 2020. (Cit   aux pages 51 and 94.)
- [98] J. SINGH, R. ZEIER, T. CALARCO et F. MOTZOI, « Compensating for nonlinear distortions in controlled quantum systems », *Phys. Rev. Appl.*, vol. 19, p. 064067, Jun 2023. (Cit      la page 51.)
- [99] S. LAZ  R, Q. FICHEUX, J. HERRMANN, A. REMM, N. LACROIX, C. HELLINGS, F. SWIADEK, D. C. ZANUZ, G. J. NORRIS, M. B. PANAH, A. FLASBY, M. KERSCHBAUM, J.-C. BESSE, C. EICHLER et A. WALLRAFF, « Calibration of drive nonlinearity for arbitrary-angle single-qubit gates using error amplification », *Phys. Rev. Appl.*, vol. 20, p. 024036, Aug 2023. (Cit      la page 51.)
- [100] E. HYYP  , A. VEPS  L  INEN, M. PAPI  , C. F. CHAN, S. INEL, A. LANDRA, W. LIU, J. LUUS, F. MARXER, C. OCKELOEN-KORPPI, S. ORBELL, B. TARASINSKI et J. HEIN-

- soo, « Reducing leakage of single-qubit gates for superconducting quantum processors using analytical control pulse envelopes », 2024. arXiv :2402.17757. (Cité à la page 51.)
- [101] B. FOXEN, C. NEILL, A. DUNSWORTH, P. ROUSHAN, B. CHIARO, A. MEGRANT, J. KELLY, Z. CHEN, K. SATZINGER, R. BARENDTS, F. ARUTE, K. ARYA, R. BABBUSH, D. BACON, J. BARDIN, S. BOIXO, D. BUELL, B. BURKETT, Y. CHEN *et al.*, « Demonstrating a continuous set of two-qubit gates for near-term quantum algorithms », *Physical Review Letters*, vol. 125, sept. 2020. (Cité aux pages 51, 64, 88, and 97.)
- [102] É. GENOIS, J. A. GROSS, A. DI PAOLO, N. J. STEVENSON, G. KOOLSTRA, A. HASHIM, I. SIDDIQI et A. BLAIS, « Quantum-tailored machine-learning characterization of a superconducting qubit », *PRX Quantum*, vol. 2, p. 040355, Dec 2021. (Cité aux pages 52, 53, and 58.)
- [103] E. FLURIN, L. S. MARTIN, S. HACOEN-GOURGY et I. SIDDIQI, « Using a Recurrent Neural Network to Reconstruct Quantum Dynamics of a Superconducting Qubit from Physical Observations », *Physical Review X*, vol. 10, p. 011006, jan. 2020. (Cité aux pages 52 and 58.)
- [104] A. P. PEIRCE, M. A. DAHLEH et H. RABITZ, « Optimal control of quantum-mechanical systems : Existence, numerical approximation, and applications », *Phys. Rev. A*, vol. 37, p. 4950–4964, Jun 1988. (Cité à la page 54.)
- [105] J. WERSCHNIK et E. GROSS, « Quantum optimal control theory », *Journal of Physics B : Atomic, Molecular and Optical Physics*, vol. 40, no. 18, p. R175, 2007. (Cité à la page 54.)
- [106] L. PONTRYAGIN, V. BOLTYANSKII, R. GAMKRELIDZE et E. MISHECHENKO, *The Mathematical Theory of Optimal Processes*. CRC Press, 01 1962. (Cité aux pages 54, 144, and 145.)
- [107] E. D. SONTAG, *Mathematical Control Theory*. Springer New York, NY, 07 1998. (Cité à la page 54.)
- [108] D. D’ALESSANDRO, *Introduction to quantum control and dynamics*. Chapman & Hall/CRC applied mathematics & nonlinear science, Hoboken, NJ : Taylor & Francis Ltd, 2007. (Cité à la page 54.)
- [109] U. BOSCAIN, M. SIGALOTTI et D. SUGNY, « Introduction to the pontryagin maximum principle for quantum optimal control », *PRX Quantum*, vol. 2, p. 030203, Sep 2021. (Cité à la page 54.)
- [110] G. DIRR et U. HELMKE, *Lie Theory for Quantum Control*. Wiley, 2008. (Cité à la page 54.)
- [111] N. KHANEJA, T. REISS, C. KEHLET, T. SCHULTE-HERBRÜGGEN et S. J. GLASER, « Optimal control of coupled spin dynamics : design of NMR pulse sequences by gradient ascent algorithms », *Journal of Magnetic Resonance*, vol. 172, p. 296–305, fév. 2005. (Cité aux pages 54, 55, and 144.)
- [112] P. de FOUQUIERES, S. SCHIRMER, S. GLASER et I. KUPROV, « Second order gradient ascent pulse engineering », *Journal of Magnetic Resonance*, vol. 212, p. 412–417, oct. 2011. (Cité à la page 56.)
- [113] S. MACHNES, U. SANDER, S. J. GLASER, P. de FOUQUIÈRES, A. GRUSLYS, S. SCHIRMER et T. SCHULTE-HERBRÜGGEN, « Comparing, optimizing, and benchmarking quantum-control algorithms in a unifying programming framework », *Phys. Rev. A*, vol. 84, p. 022305, Aug 2011. (Cité à la page 56.)

- [114] V. KROTOV, *Global Methods in Optimal Control Theory*. CRC Press, oct. 1995. (Cité à la page 56.)
- [115] D. M. REICH, M. NDONG et C. P. KOCH, « Monotonically convergent optimization in quantum control using krotov's method », *The Journal of Chemical Physics*, vol. 136, mars 2012. (Cité à la page 56.)
- [116] D. J. EGGER et F. K. WILHELM, « Adaptive hybrid optimal quantum control for imprecisely characterized systems », *Phys. Rev. Lett.*, vol. 112, p. 240503, Jun 2014. (Cité aux pages 56, 97, and 140.)
- [117] J. KELLY, R. BARENDS, B. CAMPBELL, Y. CHEN, Z. CHEN, B. CHIARO, A. DUNSWORTH, A. G. FOWLER, I.-C. HOI, E. JEFFREY, A. MEGRANT, J. MUTUS, C. NEILL, P. J. J. O'MALLEY, C. QUINTANA, P. ROUSHAN, D. SANK, A. VAINSENCHER, J. WENNER, T. C. WHITE, A. N. CLELAND et J. M. MARTINIS, « Optimal quantum control using randomized benchmarking », *Phys. Rev. Lett.*, vol. 112, p. 240504, Jun 2014. (Cité aux pages 56, 97, and 140.)
- [118] M. WERNINGHAUS, D. J. EGGER, F. ROY, S. MACHNES, F. K. WILHELM et S. FILIPP, « Leakage reduction in fast superconducting qubit gates via optimal control », *npj Quantum Information*, vol. 7, jan. 2021. (Cité aux pages 56 and 140.)
- [119] J. J. BURNETT, A. BENGTSOON, M. SCIGLIUZZO, D. NIEPCE, M. KUDRA, P. DELSING et J. BYLANDER, « Decoherence benchmarking of superconducting qubits », *npj Quantum Information*, vol. 5, juin 2019. (Cité à la page 56.)
- [120] S. SCHLÖR, J. LISENFELD, C. MÜLLER, A. BILMES, A. SCHNEIDER, D. P. PAPPAS, A. V. USTINOV et M. WEIDES, « Correlating decoherence in transmon qubits : Low frequency noise by single fluctuators », *Phys. Rev. Lett.*, vol. 123, p. 190502, Nov 2019. (Cité à la page 56.)
- [121] OPENAI, « Gpt-4 technical report », 2024. arXiv :2303.08774. (Cité à la page 57.)
- [122] J. BIAMONTE, P. WITTEK, N. PANCOTTI, P. REBENTROST, N. WIEBE et S. LLOYD, « Quantum machine learning », *Nature*, vol. 549, p. 195–202, sept. 2017. (Cité à la page 57.)
- [123] V. DUNJKO et H. J. BRIEGEL, « Machine learning & artificial intelligence in the quantum domain : a review of recent progress », *Reports on Progress in Physics*, vol. 81, p. 074001, jun 2018. (Cité à la page 57.)
- [124] P. MEHTA, M. BUKOV, C.-H. WANG, A. G. DAY, C. RICHARDSON, C. K. FISHER et D. J. SCHWAB, « A high-bias, low-variance introduction to machine learning for physicists », *Physics Reports*, vol. 810, p. 1–124, mai 2019. (Cité à la page 57.)
- [125] G. CARLEO, I. CIRAC, K. CRANMER, L. DAUDET, M. SCHULD, N. TISHBY, L. VOGT-MARANTO et L. ZDEBOROVÁ, « Machine learning and the physical sciences », *Reviews of Modern Physics*, vol. 91, déc. 2019. (Cité à la page 57.)
- [126] F. MARQUARDT, « Machine learning and quantum devices », *SciPost Physics Lecture Notes*, mai 2021. (Cité à la page 57.)
- [127] I. GOODFELLOW, Y. BENGIO et A. COURVILLE, *Deep Learning*. MIT Press, 2016. <http://www.deeplearningbook.org>. (Cité aux pages 57 and 58.)
- [128] Y. LECUN, Y. BENGIO et G. HINTON, « Deep learning », *Nature*, vol. 521, p. 436–444, mai 2015. (Cité à la page 57.)

- [129] R. S. SUTTON et A. G. BARTO, *Reinforcement Learning : An Introduction*. The MIT Press, second éd., 2018. (Cité aux pages 57 and 60.)
- [130] D. FOSTER et K. FRISTON, *Generative Deep Learning : Teaching Machines to Paint, Write, Compose, and Play*. O'Reilly, 2023. (Cité à la page 57.)
- [131] K. FUKUSHIMA, « Neocognitron : A self-organizing neural network model for a mechanism of pattern recognition unaffected by shift in position », *Biological Cybernetics*, vol. 36, p. 193–202, avril 1980. (Cité à la page 58.)
- [132] Y. LECUN, B. BOSER, J. DENKER, D. HENDERSON, R. HOWARD, W. HUBBARD et L. JACKEL, « Handwritten Digit Recognition with a Back-Propagation Network », in *Advances in Neural Information Processing Systems*, vol. 2, Morgan-Kaufmann, 1989. (Cité à la page 58.)
- [133] S. LAWRENCE, C. GILES, A. C. TSOI et A. BACK, « Face recognition : a convolutional neural-network approach », *IEEE Transactions on Neural Networks*, vol. 8, p. 98–113, jan. 1997. (Cité à la page 58.)
- [134] R. J. WILLIAMS et D. ZIPSER, « A learning algorithm for continually running fully recurrent neural networks », *Neural Computation*, vol. 1, no. 2, p. 270–280, 1989. (Cité à la page 58.)
- [135] M. HIBAT-ALLAH, M. GANAHL, L. E. HAYWARD, R. G. MELKO et J. CARRASQUILLA, « Recurrent neural network wave functions », *Physical Review Research*, vol. 2, juin 2020. (Cité à la page 58.)
- [136] M. AUGUST et X. NI, « Using recurrent neural networks to optimize dynamical decoupling for quantum memory », *Physical Review A*, vol. 95, jan. 2017. (Cité à la page 58.)
- [137] L. BANCHI, E. GRANT, A. ROCCHETTO et S. SEVERINI, « Modelling non-markovian quantum processes with recurrent neural networks », *New Journal of Physics*, vol. 20, p. 123030, déc. 2018. (Cité à la page 58.)
- [138] S. HOCHREITER et J. SCHMIDHUBER, « Long Short-Term Memory », *Neural Computation*, vol. 9, p. 1735–1780, 11 1997. (Cité à la page 58.)
- [139] K. HORNIK, M. STINCHCOMBE et H. WHITE, « Multilayer feedforward networks are universal approximators », *Neural Networks*, vol. 2, no. 5, p. 359–366, 1989. (Cité à la page 59.)
- [140] R. ORÚS, « Tensor networks for complex quantum systems », *Nature Reviews Physics*, vol. 1, p. 538–550, août 2019. (Cité à la page 59.)
- [141] M. CRAMER, M. B. PLENIO, S. T. FLAMMIA, R. SOMMA, D. GROSS, S. D. BARTLETT, O. LANDON-CARDINAL, D. POULIN et Y.-K. LIU, « Efficient quantum state tomography », *Nature Communications*, vol. 1, déc. 2010. (Cité à la page 60.)
- [142] C. JAMESON, Z. QIN, A. GOLDAR, M. B. WAKIN, Z. ZHU et Z. GONG, « Optimal quantum state tomography with local informationally complete measurements », 2024. arXiv :2408.07115. (Cité à la page 60.)
- [143] G. TORLAI, C. J. WOOD, A. ACHARYA, G. CARLEO, J. CARRASQUILLA et L. AOLITA, « Quantum process tomography with unsupervised learning and tensor networks », *Nature Communications*, vol. 14, mai 2023. (Cité à la page 60.)
- [144] C. E. GRANADE, C. FERRIE, N. WIEBE et D. G. CORY, « Robust online hamiltonian learning », *New Journal of Physics*, vol. 14, p. 103013, oct. 2012. (Cité à la page 60.)

- [145] H.-Y. HUANG, Y. TONG, D. FANG et Y. SU, « Learning many-body hamiltonians with heisenberg-limited scaling », *Phys. Rev. Lett.*, vol. 130, p. 200403, May 2023. (Cit      la page 60.)
- [146] M. BUKOV, A. G.-R. DAY, D. SELS, P. WEINBERG, A. POLKOVNIKOV et P. MEHTA, « Reinforcement Learning in Different Phases of Quantum Control », *Physical Review X*, vol. 8, p. 031086, sept. 2018. (Cit   aux pages 60 and 61.)
- [147] S. BORAH, B. SARMA, M. KEWMING, G. J. MILBURN et J. TWAMLEY, « Measurement-based feedback quantum control with deep reinforcement learning for a double-well nonlinear potential », *Phys. Rev. Lett.*, vol. 127, p. 190403, Nov 2021. (Cit      la page 60.)
- [148] V. V. SIVAK, A. EICKBUSCH, H. LIU, B. ROYER, I. TSIOUTSIOS et M. H. DEVORET, « Model-free quantum control with reinforcement learning », *Phys. Rev. X*, vol. 12, p. 011059, Mar 2022. (Cit   aux pages 60, 61, and 62.)
- [149] R. POROTTI, A. ESSIG, B. HUARD et F. MARQUARDT, « Deep reinforcement learning for quantum state preparation with weak nonlinear measurements », *Quantum*, vol. 6, p. 747, juin 2022. (Cit   aux pages 60 and 61.)
- [150] R. POROTTI, V. PEANO et F. MARQUARDT, « Gradient-ascent pulse engineering with feedback », *PRX Quantum*, vol. 4, p. 030305, Jul 2023. (Cit   aux pages 60 and 90.)
- [151] L. DING, M. HAYS, Y. SUNG, B. KANNAN, J. AN, A. DI PAOLO, A. H. KARAMLOU, T. M. HAZARD, K. AZAR, D. K. KIM, B. M. NIEDZIELSKI, A. MELVILLE, M. E. SCHWARTZ, J. L. YODER, T. P. ORLANDO, S. GUSTAVSSON, J. A. GROVER, K. SERNIK et W. D. OLIVER, « High-fidelity, frequency-flexible two-qubit fluxonium gates with a transmon coupler », *Phys. Rev. X*, vol. 13, p. 031035, Sep 2023. (Cit   aux pages 60, 62, and 140.)
- [152] C. J. C. H. WATKINS et P. DAYAN, « Q-learning », *Machine Learning*, vol. 8, p. 279–292, mai 1992. (Cit      la page 61.)
- [153] J. SCHULMAN, F. WOLSKI, P. DHARIWAL, A. RADFORD et O. KLIMOV, « Proximal policy optimization algorithms », 2017. arXiv :1707.06347. (Cit      la page 61.)
- [154] B. VLASTAKIS, G. KIRCHMAIR, Z. LEGHTAS, S. E. NIGG, L. FRUNZIO, S. M. GIRVIN, M. MIRRAHIMI, M. H. DEVORET et R. J. SCHOELKOPF, « Deterministically encoding quantum information using 100-photon schr  dinger cat states », *Science*, vol. 342, no. 6158, p. 607–610, 2013. (Cit      la page 62.)
- [155] A. ALMANAKLY, B. YANKELEVICH, M. HAYS, B. KANNAN, R. ASSOULY, A. GREENE, M. GINGRAS, B. M. NIEDZIELSKI, H. STICKLER, M. E. SCHWARTZ, K. SERNIK, J. I.-J. WANG, T. P. ORLANDO, S. GUSTAVSSON, J. A. GROVER et W. D. OLIVER, « Deterministic remote entanglement using a chiral quantum interconnect », 2024. arXiv :2408.05164. (Cit      la page 62.)
- [156] A. GRIEWANK et A. WALTHER, *Evaluating Derivatives*. Other Titles in Applied Mathematics, Society for Industrial and Applied Mathematics, jan. 2008. (Cit      la page 63.)
- [157] A. G. BAYDIN, B. A. PEARLMUTTER, A. A. RADUL et J. M. SISKIND, « Automatic differentiation in machine learning : a survey », 2018. arXiv :1502.05767. (Cit   aux pages 63 and 144.)
- [158] F. PRETI, T. CALARCO et F. MOTZOI, « Continuous quantum gate sets and pulse-class meta-optimization », *PRX Quantum*, vol. 3, oct. 2022. (Cit      la page 64.)
- [159] A. PASZKE, S. GROSS, F. MASSA, A. LERER, J. BRADBURY, G. CHANAN, T. KILLEEN, Z. LIN, N. GIMELSHEIN, L. ANTIGA, A. DESMAISON, A. KOPF, E. YANG, Z. DEVITO,

- M. RAISON, A. TEJANI, S. CHILAMKURTHY, B. STEINER, L. FANG, J. BAI et S. CHINTALA, « Pytorch : An imperative style, high-performance deep learning library », in *Advances in Neural Information Processing Systems*, vol. 32, 2019. (Cité aux pages 65 and 151.)
- [160] A. YOUSSEY, Y. YANG, R. J. CHAPMAN, B. HAYLOCK, F. LENZINI, M. LOBINO et A. PERUZZO, « Experimental graybox quantum system identification and control », *npj Quantum Information*, vol. 10, p. 1–9, jan. 2024. (Cité à la page 85.)
- [161] T. HASTIE, R. TIBSHIRANI et J. FRIEDMAN, *The Elements of Statistical Learning*. Springer New York, août 2009. (Cité aux pages 86 and 140.)
- [162] A. HASHIM, R. K. NAIK, A. MORVAN, J.-L. VILLE, B. MITCHELL, J. M. KREIKEBAUM, M. DAVIS, E. SMITH, C. IANCU, K. P. O'BRIEN, I. HINCKS, J. J. WALLMAN, J. EMERSON et I. SIDDIQI, « Randomized compiling for scalable quantum computing on a noisy superconducting quantum processor », *Phys. Rev. X*, vol. 11, p. 041039, Nov 2021. (Cité aux pages 86 and 96.)
- [163] J. B. HERTZBERG, E. J. ZHANG, S. ROSENBLATT, E. MAGESAN, J. A. SMOLIN, J.-B. YAU, V. P. ADIGA, M. SANDBERG, M. BRINK, J. M. CHOW et J. S. ORCUTT, « Laser-annealing josephson junctions for yielding scaled-up superconducting quantum processors », *npj Quantum Information*, vol. 7, août 2021. (Cité à la page 87.)
- [164] F. YAN, P. KRANTZ, Y. SUNG, M. KJAERGAARD, D. L. CAMPBELL, T. P. ORLANDO, S. GUSTAVSSON et W. D. OLIVER, « Tunable coupling scheme for implementing high-fidelity two-qubit gates », *Physical Review Applied*, vol. 10, nov. 2018. (Cité à la page 87.)
- [165] R. SHILLITO, J. A. GROSS, A. DI PAOLO, É. GENOIS et A. BLAIS, « Fast and differentiable simulation of driven quantum systems », *Physical Review Research*, vol. 3, p. 033266, sept. 2021. (Cité à la page 88.)
- [166] N. SUNDARESAN, I. LAUER, E. PRITCHETT, E. MAGESAN, P. JURCEVIC et J. M. GAMBETTA, « Reducing unitary and spectator errors in cross resonance with optimized rotary echoes », *PRX Quantum*, vol. 1, p. 020318, Dec 2020. (Cité à la page 88.)
- [167] P. GUILMIN, P. ROUCHON et A. TILLOY, « Correlation functions for realistic continuous quantum measurement », *IFAC-PapersOnLine*, vol. 56, no. 2, p. 5164–5170, 2023. 22nd IFAC World Congress. (Cité à la page 89.)
- [168] A. S. ROY, K. PACK, N. WITTLER et S. MACHNES, « Software tool-set for automated quantum system identification and device bring up », 2022. arXiv :2205.04829. (Cité à la page 89.)
- [169] T. M. MOERLAND, J. BROEKENS, A. PLAAT et C. M. JONKER, « Model-based reinforcement learning : A survey », 2022. arXiv :2006.16712. (Cité à la page 90.)
- [170] F. ARUTE, K. ARYA, R. BABBUSH, D. BACON, J. C. BARDIN, R. BARENDTS, A. BENGTSSON, S. BOIXO, M. BROUGHTON, B. B. BUCKLEY, D. A. BUELL, B. BURKETT, N. BUSHNELL, Y. CHEN, Z. CHEN *et al.*, « Observation of separated dynamics of charge and spin in the Fermi-Hubbard model », *arXiv :2010.07965*, oct. 2020. (Cité aux pages 93 and 99.)
- [171] N. MOLL, P. BARKOUTSOS, L. S. BISHOP, J. M. CHOW, A. CROSS, D. J. EGGER, S. FILIPP, A. FUHRER, J. M. GAMBETTA, M. GANZHORN, A. KANDALA, A. MEZZACAPO, P. MÜLLER, W. RIESS, G. SALIS, J. SMOLIN, I. TAVERNELLI et K. TEMME, « Quantum

- optimization using variational algorithms on near-term quantum devices », *Quantum Science and Technology*, vol. 3, p. 030503, juin 2018. (Cité à la page 93.)
- [172] F. ARUTE, K. ARYA, R. BABBUSH, D. BACON, J. C. BARDIN, R. BARENDTS, R. BISWAS, S. BOIXO, F. G. S. L. BRANDAO, D. A. BUELL, B. BURKETT, Y. CHEN, Z. CHEN, B. CHIARO, R. COLLINS, W. COURTNEY, A. DUNSWORTH, E. FARHI, B. FOXEN, A. FOWLER, C. GIDNEY *et al.*, « Quantum supremacy using a programmable superconducting processor », *Nature*, vol. 574, p. 505–510, oct. 2019. (Cité aux pages 93, 94, 95, 96, and 102.)
- [173] C. C. BULTINK, M. A. ROL, T. E. O'BRIEN, X. FU, B. C. S. DIKKEN, C. DICKEL, R. F. L. VERMEULEN, J. C. de STERKE, A. BRUNO, R. N. SCHOUTEN *et* L. DiCARLO, « Active resonator reset in the nonlinear dispersive regime of circuit qed », *Phys. Rev. Appl.*, vol. 6, p. 034008, Sep 2016. (Cité à la page 94.)
- [174] L. VIOLA, E. KNILL *et* S. LLOYD, « Dynamical decoupling of open quantum systems », *Physical Review Letters*, vol. 82, p. 2417–2421, mars 1999. (Cité aux pages 95 and 143.)
- [175] J. EMERSON, R. ALICKI *et* K. ŻYCZKOWSKI, « Scalable noise estimation with random unitary operators », *Journal of Optics B : Quantum and Semiclassical Optics*, vol. 7, p. S347–S352, sept. 2005. (Cité à la page 95.)
- [176] S. BOIXO, S. V. ISAKOV, V. N. SMELYANSKIY, R. BABBUSH, N. DING, Z. JIANG, M. J. BREMNER, J. M. MARTINIS *et* H. NEVEN, « Characterizing quantum supremacy in near-term devices », *Nature Physics*, vol. 14, p. 595–600, avril 2018. (Cité à la page 95.)
- [177] J. EISERT, D. HANGLEITER, N. WALK, I. ROTH, D. MARKHAM, R. PAREKH, U. CHAUBAUD *et* E. KASHEFI, « Quantum certification and benchmarking », *Nature Reviews Physics*, vol. 2, p. 382–390, juin 2020. (Cité à la page 96.)
- [178] J. WALLMAN, C. GRANADE, R. HARPER *et* S. T. FLAMMIA, « Estimating the coherence of noise », *New Journal of Physics*, vol. 17, p. 113020, nov. 2015. (Cité à la page 96.)
- [179] M. McEWEN, D. BACON *et* C. GIDNEY, « Relaxing hardware requirements for surface code circuits using time-dynamics », *Quantum*, vol. 7, p. 1172, nov. 2023. (Cité à la page 98.)
- [180] D. C. MCKAY, C. J. WOOD, S. SHELDON, J. M. CHOW *et* J. M. GAMBETTA, « Efficient Z gates for quantum computing », *Physical Review A*, vol. 96, août 2017. (Cité à la page 98.)
- [181] G. FLOQUET, « Sur les équations différentielles linéaires à coefficients périodiques », *Annales scientifiques de l'École Normale Supérieure*, vol. 12, p. 47–88, 1883. (Cité à la page 99.)
- [182] A. MAUDSLEY, « Modified carr-purcell-meiboom-gill sequence for nmr fourier imaging applications », *Journal of Magnetic Resonance (1969)*, vol. 69, no. 3, p. 488–491, 1986. (Cité à la page 100.)
- [183] J. M. MARTINIS, S. NAM, J. AUMENTADO, K. M. LANG *et* C. URBINA, « Decoherence of a superconducting qubit due to bias noise », *Phys. Rev. B*, vol. 67, p. 094510, Mar 2003. (Cité à la page 103.)
- [184] Á. MÁRTON *et* J. K. ASBÓTH, « Coherent errors and readout errors in the surface code », *Quantum*, vol. 7, p. 1116, sept. 2023. (Cité à la page 141.)

- [185] P. M. HARRINGTON, E. J. MUELLER et K. W. MURCH, « Engineered dissipation for quantum information science », *Nature Reviews Physics*, vol. 4, p. 660–671, août 2022. (Cité aux pages 143 and 172.)
- [186] C. P. KOCH, « Controlling open quantum systems : tools, achievements, and limitations », *Journal of Physics : Condensed Matter*, vol. 28, p. 213001, mai 2016. (Cité à la page 143.)
- [187] D. A. LIDAR, I. L. CHUANG et K. B. WHALEY, « Decoherence-free subspaces for quantum computation », *Phys. Rev. Lett.*, vol. 81, p. 2594–2597, Sep 1998. (Cité à la page 143.)
- [188] S. MACHNES, E. ASSÉMAT, D. TANNOR et F. K. WILHELM, « Tunable, flexible, and efficient optimization of control pulses for practical qubits », *Phys. Rev. Lett.*, vol. 120, p. 150401, Apr 2018. (Cité à la page 144.)
- [189] L. HAN et M. NEUMANN, « Effect of dimensionality on the nelder–mead simplex method », *Optimization Methods and Software*, vol. 21, no. 1, p. 1–16, 2006. (Cité à la page 144.)
- [190] M. GOERZ, D. BASILEWITSCH, F. GAGO-ENCINAS, M. G. KRAUSS, K. P. HORN, D. M. REICH et C. KOCH, « Krotov : A python implementation of krotov’s method for quantum optimal control », *SciPost physics*, vol. 7, no. 6, p. 080, 2019. (Cité à la page 144.)
- [191] D. BASILEWITSCH, C. P. KOCH et D. M. REICH, « Quantum optimal control for mixed state squeezing in cavity optomechanics », *Advanced Quantum Technologies*, vol. 2, no. 3-4, p. 1800110, 2019. (Cité à la page 144.)
- [192] T. SCHULTE-HERBRÜGGEN, A. SPÖRL, N. KHANEJA et S. GLASER, « Optimal control for generating quantum gates in open dissipative systems », *Journal of Physics B : Atomic, Molecular and Optical Physics*, vol. 44, no. 15, p. 154013, 2011. (Cité à la page 144.)
- [193] S. BOUTIN, C. K. ANDERSEN, J. VENKATRAMAN, A. J. FERRIS et A. BLAIS, « Resonator reset in circuit qed by optimal control for large open quantum systems », *Phys. Rev. A*, vol. 96, p. 042315, Oct 2017. (Cité à la page 144.)
- [194] N. LEUNG, M. ABDELHAFEZ, J. KOCH et D. SCHUSTER, « Speedup for quantum optimal control from automatic differentiation based on graphics processing units », *Phys. Rev. A*, vol. 95, p. 042318, Apr 2017. (Cité aux pages 145 and 148.)
- [195] M. ABDELHAFEZ, D. I. SCHUSTER et J. KOCH, « Gradient-based optimal control of open quantum systems using quantum trajectories and automatic differentiation », *Phys. Rev. A*, vol. 99, p. 052327, May 2019. (Cité à la page 145.)
- [196] R. T. Q. CHEN, Y. RUBANOVA, J. BETTENCOURT et D. DUVENAUD, « Neural ordinary differential equations », 2019. arXiv :1806.07366. (Cité aux pages 145 and 151.)
- [197] P. KIDGER, « On neural differential equations », 2022. arXiv :2202.02435. (Cité aux pages 146 and 168.)
- [198] M. F. DUMAS, B. GROLEAU-PARÉ, A. McDONALD, M. H. MUÑOZ-ARIAS, C. LLEDÓ, B. D’ANJOU et A. BLAIS, « Unified picture of measurement-induced ionization in the transmon », 2024. arXiv :2402.06615. (Cité à la page 146.)
- [199] Y. LU, S. JOSHI, V. SAN DINH et J. KOCH, « Optimal control of large quantum systems : assessing memory and runtime performance of grape », *Journal of Physics Communications*, vol. 8, p. 025002, fév. 2024. (Cité à la page 148.)

- [200] F. SWIADEK, R. SHILLITO, P. MAGNARD, A. REMM, C. HELLINGS, N. LACROIX, Q. FICHEUX, D. C. ZANUZ, G. J. NORRIS, A. BLAIS, S. KRINNER et A. WALLRAFF, « Enhancing dispersive readout of superconducting qubits through dynamic control of the dispersive shift : Experiment and theory », 2023. arXiv :2307.07765. (Cit      la page 148.)
- [201] J. IKONEN, J. GOETZ, J. ILVES, A. KER  NEN, A. M. GUNYHO, M. PARTANEN, K. Y. TAN, D. HAZRA, L. GR  NBERG, V. VESTERINEN, S. SIMBIEROWICZ, J. HASSEL et M. M  TT  NEN, « Qubit measurement by multichannel driving », *Phys. Rev. Lett.*, vol. 122, p. 080503, Feb 2019. (Cit      la page 149.)
- [202] S. TOUZARD, A. KOU, N. E. FRATTINI, V. V. SIVAK, S. PURI, A. GRIMM, L. FRUNZIO, S. SHANKAR et M. H. DEVORET, « Gated conditional displacement readout of superconducting qubits », *Phys. Rev. Lett.*, vol. 122, p. 080502, Feb 2019. (Cit      la page 149.)
- [203] N. DIDIER, J. BOURASSA et A. BLAIS, « Fast quantum nondemolition readout by parametric modulation of longitudinal qubit-oscillator interaction », *Phys. Rev. Lett.*, vol. 115, p. 203601, Nov 2015. (Cit      la page 149.)
- [204] M. H. MU  OZ ARIAS, C. LLED   et A. BLAIS, « Qubit readout enabled by qubit cloaking », *Phys. Rev. Appl.*, vol. 20, p. 054013, Nov 2023. (Cit      la page 149.)
- [205] F. MALLET, F. R. ONG, A. PALACIOS-LALOY, F. NGUYEN, P. BERTET, D. VION et D. ESTEVE, « Single-shot qubit readout in circuit quantum electrodynamics », *Nat Phys*, vol. 5, no. 11, p. 791–795, 2009. (Cit      la page 149.)
- [206] M. PECHAL, L. HUTHMACHER, C. EICHLER, S. ZEYTINO  LU, A. A. ABDUMALIKOV, S. BERGER, A. WALLRAFF et S. FILIPP, « Microwave-controlled generation of shaped single photons in circuit quantum electrodynamics », *Phys. Rev. X*, vol. 4, p. 041010, Oct 2014. (Cit      la page 150.)
- [207] D. P. KINGMA et J. BA, « Adam : A method for stochastic optimization », 2017. arXiv :1412.6980. (Cit      la page 170.)
- [208] W. CAI, Y. MA, W. WANG, C.-L. ZOU et L. SUN, « Bosonic quantum error correction codes in superconducting quantum circuits », *Fundamental Research*, vol. 1, p. 50–67, jan. 2021. (Cit   aux pages 169 and 172.)
- [209] D. GOTTESMAN, A. KITAEV et J. PRESKILL, « Encoding a qubit in an oscillator », *Phys. Rev. A*, vol. 64, p. 012310, Jun 2001. (Cit      la page 169.)
- [210] B. ROYER, S. SINGH et S. M. GIRVIN, « Stabilization of finite-energy gottesman-kitaev-preskill states », *Phys. Rev. Lett.*, vol. 125, p. 260509, Dec 2020. (Cit      la page 170.)
- [211] B. de NEEVE, T. L. NGUYEN, T. BEHRLE et J. HOME, « Error correction of a logical grid state qubit by dissipative pumping », 2020. arXiv :2010.09681. (Cit      la page 170.)
- [212] V. V. ALBERT, K. NOH, K. DUVENVOORDEN, D. J. YOUNG, R. T. BRIERLEY, P. REINHOLD, C. VUILLOT, L. LI, C. SHEN, S. M. GIRVIN, B. M. TERHAL et L. JIANG, « Performance and structure of single-mode bosonic codes », *Phys. Rev. A*, vol. 97, p. 032346, Mar 2018. (Cit      la page 170.)
- [213] K. NOH, V. V. ALBERT et L. JIANG, « Quantum capacity bounds of gaussian thermal loss channels and achievable rates with gottesman-kitaev-preskill codes », *IEEE Transactions on Information Theory*, vol. 65, no. 4, p. 2563–2582, 2019. (Cit      la page 170.)

- [214] C. DEGEN, F. REINHARD et P. CAPPELLARO, « Quantum sensing », *Reviews of Modern Physics*, vol. 89, juil. 2017. (Cité à la page 172.)
- [215] K. AZUMA, S. E. ECONOMOU, D. ELKOSS, P. HILAIRE, L. JIANG, H.-K. LO et I. TZITRIN, « Quantum repeaters : From quantum networks to the quantum internet », *Rev. Mod. Phys.*, vol. 95, p. 045006, Dec 2023. (Cité à la page 172.)
- [216] A. MORVAN, B. VILLALONGA, X. MI, S. MANDRÀ, A. BENGTSSON, P. V. KLIMOV, Z. CHEN, S. HONG, C. ERICKSON, I. K. DROZDOV, J. CHAU, G. LAUN, R. MOVASSAGH, A. ASFAW, L. T. A. N. BRANDÃO, R. PERALTA, D. ABANIN, R. ACHARYA, R. ALLEN *et al.*, « Phase transition in random circuit sampling », 2023. arXiv :2304.11119. (Cité à la page 172.)
- [217] X. MI, A. A. MICHAILIDIS, S. SHABANI, K. C. MIAO, P. V. KLIMOV, J. LLOYD, E. ROSENBERG, R. ACHARYA, I. ALEINER, T. I. ANDERSEN, M. ANSMANN, F. ARUTE, K. ARYA, A. ASFAW, J. ATALAYA, J. C. BARDIN, A. BENGTSSON *et al.*, « Stable quantum-correlated many-body states through engineered dissipation », *Science*, vol. 383, p. 1332–1337, mars 2024. (Cité à la page 172.)
- [218] E. ROSENBERG, T. I. ANDERSEN, R. SAMAJDAR, A. PETUKHOV, J. C. HOKE, D. ABANIN, A. BENGTSSON, I. K. DROZDOV, C. ERICKSON, P. V. KLIMOV, X. MI, A. MORVAN, M. NEELEY, C. NEILL, R. ACHARYA, R. ALLEN *et al.*, « Dynamics of magnetization at infinite temperature in a heisenberg spin chain », *Science*, vol. 384, p. 48–53, avril 2024. (Cité à la page 172.)