

Fault-tolerant syndrome extraction and cat state preparation with fewer qubits

Prithviraj Prabhu and Ben Reichardt

University of Southern California, Los Angeles, CA 90089, USA

We reduce the extra qubits needed for two fault-tolerant quantum computing protocols: error correction, specifically syndrome bit measurement, and cat state preparation. For distance-three fault-tolerant syndrome extraction, we show an exponential reduction in qubit overhead over the previous best protocol. For a weight- w stabilizer, we demonstrate that stabilizer measurement tolerating one fault needs at most $\lceil \log_2 w \rceil + 1$ ancilla qubits. If qubits reset quickly, four ancillas suffice. We also study the preparation of entangled cat states, and prove that the overhead for distance-three fault tolerance is logarithmic in the cat state size. These results apply both to near-term experiments with a few qubits, and to the general study of the asymptotic resource requirements of syndrome measurement and state preparation.

With a flag qubits, previous methods use $O(a)$ flag patterns to identify faults. In order to use the same flag qubits more efficiently, we show how to use nearly all 2^a possible flag patterns, by constructing maximal-length paths through the a -dimensional hypercube.

1 Introduction

A critical component of quantum error correction is syndrome measurement: a set of circuits that are used to pinpoint which qubits have errors. This process of error identification is itself susceptible to noise and may fail. To make this robust, extra (ancilla) qubits can be used to identify damaging mid-circuit faults and mitigate the spread of errors. The objective of this paper is to reduce the overhead of ancilla qubits used in imparting this fault tolerance. In particular, we focus on optimizing the flag technique for distance-three fault-tolerant stabilizer measurement. We also reduce qubit overhead in distance-three fault-tolerant cat state preparation. Cat states [1] have applications in many areas of quantum computing, including communication [2], information processing [3], and error correction [4]. Besides practical applications, our results on cat state preparation are theoretically interesting since: *i*) we introduce the study of asymptotic estimates of qubit overhead for the fault-tolerant

Prithviraj Prabhu: prithvirajprab@gmail.com

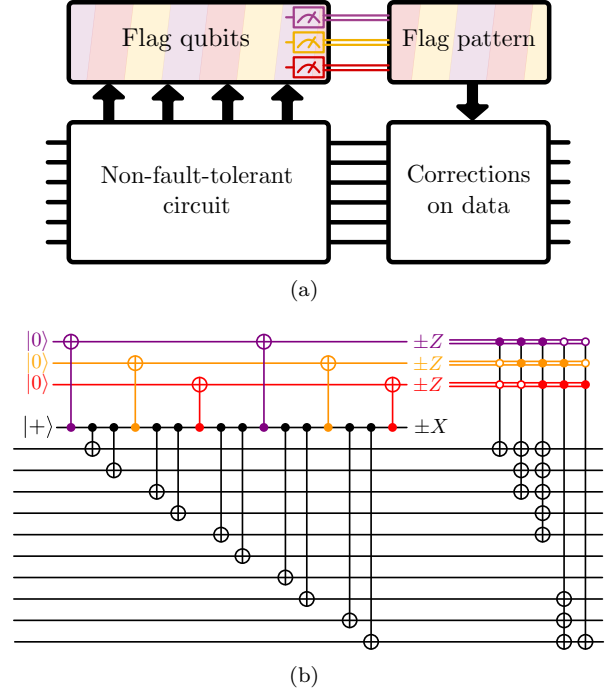


Figure 1: (a) Function of a flag scheme. Errors in a non-fault-tolerant circuit can be made to spread into flag qubits. On measurement, the flag qubits yield a pattern of 1s and 0s, based on which the data is corrected. (b) Circuit to measure the stabilizer $X^{\otimes 10}$, using three flag qubits, in color, to protect against one X fault (distance three).

preparation of cat states of *arbitrary* size, and, *ii*) ideas developed for cat state preparation may provide clues for the fault-tolerant preparation of logical states of more complex codes.

We strive for low qubit overhead since quantum computers with limited qubits count resources preciously, and even minor improvements can free up extra qubits for other tasks. In topological codes where stabilizers are localized in space and are of low weight, only a few flag qubits close to each stabilizer suffice to impart fault tolerance [7, 8, 9]. It has also been shown that with adaptive control and quickly resetting qubits, only four ancillas are required for the universal fault-tolerant operation of some distance-three codes [10, 11]. In this paper, we present a general fault-tolerant protocol that works for a stabilizer of any size. If qubits are connected well enough, we show that only logarithmic overhead is required for

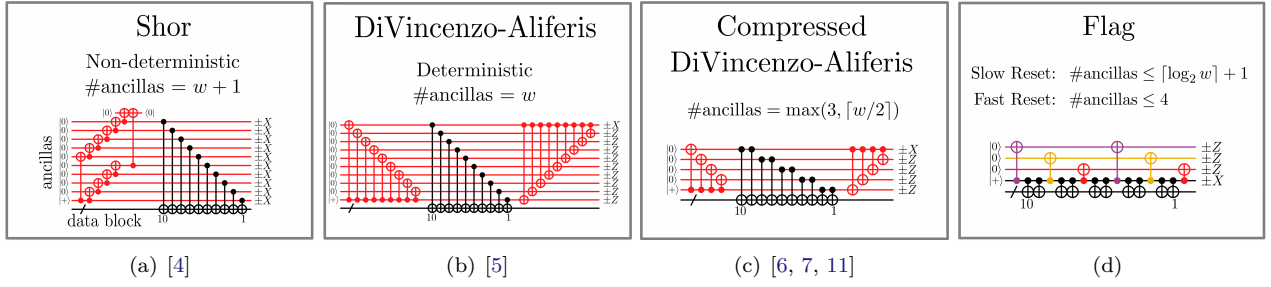


Figure 2: Historical progression of stabilizer measurement circuits, illustrated by a weight-10 X stabilizer measurement. The black CNOTs have targets on the 10 data qubits, collectively represented by a black wire. In (b-d), fault-tolerance is only guaranteed to distance three and Pauli corrections, or frame updates, are applied to the data based on the Z basis measurements. (a) Shor’s method uses $w + 1$ ancillas and requires a fault-tolerantly prepared cat state. (b,c) These methods use unverified cat states with subsequent error decoding, giving a deterministic circuit. (d) Our flag method prepares and unprepares an ancilla cat state while collecting the stabilizer. Exponentially more flag patterns can thus be accessed for fault diagnosis.

fault-tolerant stabilizer measurement, an exponential space improvement over the previous linear overhead.

The general model of flag-based fault tolerance is displayed in Fig. 1(a). Here a set of flag ancilla qubits monitor operations in a non-fault-tolerant circuit and when measured at the end, produce flag patterns that identify mid-circuit faults. Based on the observed flag pattern, a correction is applied to the data to minimize the spread of errors. As an example, Fig. 1(b) measures a stabilizer on 10 data qubits while tolerating one fault. The three colored qubits are the flags and the measured flag patterns each imply different corrections. Also note that the sequence of flag patterns 100, 110, 111, 011, 001 is a path on the hypercube and corresponds to the order of the flag CNOTs, e.g., between 100 and 110 a CNOT targets flag qubit 2.

In this paper, we restrict discussion to the measurement of individual stabilizers of a quantum code, as in Shor-style fault-tolerant stabilizer measurement [4]. We do not consider measuring multiple stabilizers in parallel, as in Refs. [12, 13, 14]. Figure 2 displays improvements made to Shor’s method. Note that Shor’s method can tolerate any number of faults by increasing the fault tolerance of the cat state preparation. The subsequent schemes forgo this property and are only fault-tolerant to distance three. DiVincenzo and Aliferis first make the circuit deterministic by removing the need for cat state verification [5]. This ensures that a circuit designer need not wait for a fault-tolerantly prepared cat state before measuring the stabilizer. Subsequent improvements were made in Refs. [6, 7, 11] to reduce ancilla count by coupling each ancilla qubit to two data qubits instead of one.

With our flag method, the ancilla cat state is prepared and unprepared while collecting the stabilizer. As in Fig. 1(b), an X fault occurring anywhere on the $|+\rangle$ qubit may spread into the data, but will also leave its imprint on the flags. This is then measured out as a flag pattern. Due to the particular arrangement of the flag CNOTs, any fault that can spread to a data error of weight more than one triggers one of the five shown flag patterns. Each flag pattern then applies a

Table 1: Cat state size for different preparation methods that use m ancilla qubit measurements.

Method	Cat state size w
<i>Deterministic Correction</i> (Theorem 5)	$w \leq 3(2^m - 2m + 2)$ depth = $(w - 1) + 2^{m-2}$
<i>Adaptive Correction</i> (Theorem 6)	$w \leq 3(2^m - 2m + 3)$
<i>Error Detection</i> (Theorem 7)	$w \leq 3 \cdot 2^{m-1}$
<i>Parallelized Correction</i> (Theorem 8)	$w = 2m = 2 \cdot 2^j, j \in \mathbb{N}$ depth = $2 + \log_2 w$

unique correction that ensures that there is at most one data qubit in error. This satisfies the condition for fault tolerance, which states that k faults in a circuit should cause no more than k qubits to have errors.

For the distance-three fault-tolerant measurement of a weight- w stabilizer, we propose two methods based on the speed of qubit reset. With fast qubit reset, Theorem 3, only three flag ancillas are required in total, but each flag needs to be measured once per four data qubits. If more flags are used in parallel, the number of accessible flag patterns grows exponentially and the number of measurements per ancilla converges to one. This is the regime of slow qubit reset, Theorem 4, which uses at most $\lceil \log_2 w \rceil$ flag ancillas measured only at the end. Additionally, we show circuits for distance-five and distance-seven fault-tolerant stabilizer measurement in Appendix C.

Table 1 contains bounds on the ancilla overhead for preparing weight- w cat states fault-tolerantly to distance-three. If the flag qubits can reset quickly, Theorem 5 states that only one flag qubit is required and it needs to be reset and measured m times. Since the flag qubits operate independently, it is also possible to use m flag qubits, with each one being measured once. We further show how to use an adaptive circuit in Theorem 6 to marginally increase the number of flag patterns in use.

Appendices A and B contain two additional circuits for distance-three weight- w cat state preparation. In Theorem 7, we show how to use postselection to prepare cat states while tolerating *two* faults. Finally Theorem 8 details how to create low-depth circuits for distance-three fault-tolerant cat state preparation, which may be useful in technologies with many qubits or long two-qubit gate times.

The rest of the paper is divided into three sections. Section 2 details the construction of the two paths on the hypercube that we use as flag sequences. Section 3 describes how to use these sequences for distance-three fault-tolerant stabilizer measurement, and Section 4 deals with cat state preparation.

2 Flag sequences

A flag pattern is a string of 1s and 0s that arises from measuring flag qubits. A flag pattern with a flags is a vertex of the a -dimensional hypercube $\{0,1\}^a$. We show how to construct two maximal-length paths through the hypercube. Between sequential flag patterns only one bit changes, which in the fault-tolerant circuit constructions below will correspond to a CNOT from the syndrome qubit to that flag qubit.

The first type of flag sequence just requires a maximal-length traversal of the a -dimensional hypercube. A simple choice is the Gray code [15, 16].

Lemma 1. *For $a \geq 1$, the Gray code creates a length- 2^a Hamming path in the a -dimensional hypercube $\{0,1\}^a$.*

Proof. We construct the sequence inductively. For $a = 1$, use 0, 1. For $a > 1$, first run the sequence for $a - 1$ with 0s appended, then run it backwards with 1s appended. $\square$

For $a = 2$, e.g., the sequence is 00, 10, 11, 01. For $a = 3$, the sequence is 000, 100, 110, 010, 011, 111, 101, 001.

The second type of sequence is related to the degree of fault tolerance of the circuit. By definition, fault tolerance to distance d implies that for all $k \leq t = \lfloor \frac{d-1}{2} \rfloor$, correlated errors of weight k occur with k -th order probability. For distance-three Calderbank-Shor-Steane (CSS) fault-tolerant syndrome measurement, any single fault should result in a data error with X and Z components having weight zero or one.

In order to ensure that the circuit is distance-three fault-tolerant, we need to ensure that a measurement fault on any one ancilla qubit does not trigger corrections of weight greater than one. Hence the second maximal-length sequence requires that there are no weight-one strings except at the start and end. As shown in Fig. 1(b), we may assign weight-one corrections to these two patterns, but for all the others, multi-qubit corrections are required.

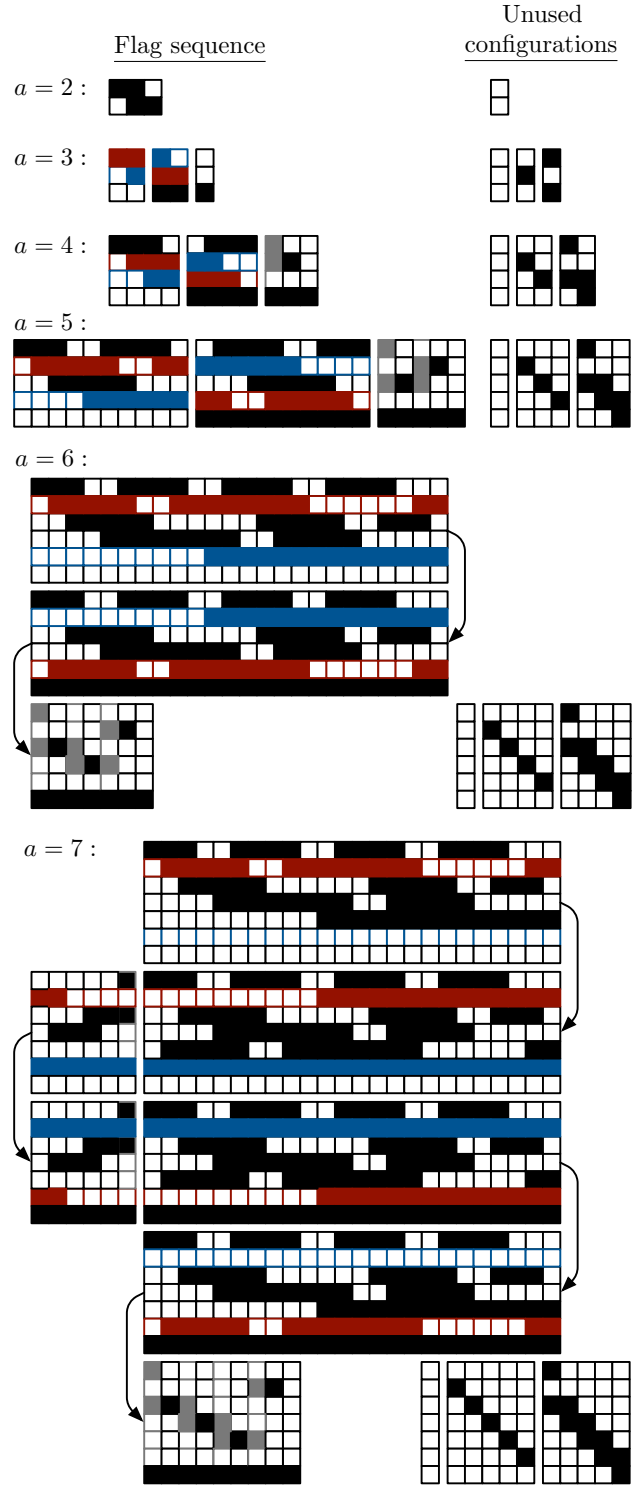


Figure 3: Flag sequences for distance-three fault-tolerant syndrome measurement, using a flag qubits, each measured once (the slow reset model). These sequences are walks through the a -dimensional hypercube, from 10^{a-1} to $0^{a-1}1$; passing through each vertex at most once and no other weight-one vertices. Flag patterns are stacked vertically and ordered initially left to right, with solid and empty squares representing 1 and 0, respectively, e.g., $\blacksquare$ represents 10, 11, 01.

Lemma 2. For $a \geq 2$, in the a -dimensional hypercube $\{0,1\}^a$ there exists a path $v_1 = 10^{a-1}, \dots, v_n = 0^{a-1}1$, with length $n = 2^a - 2a + 3$, such that all $v_2, \dots, v_{n-1}$ have weight at least two, and none repeat.

Proof. Let $1_S \in \{0,1\}^a$ denote the vertex that is 1 exactly for indices in S . Figure 3 illustrates the inductive construction. For $a = 2$, the sequence is the same as that in Lemma 1. The base case of our inductive proof is with $a = 3$, where the sequence is 100, 110, 111, 011, 001. For $a > 3$, first run the previous sequence for $b = a - 1$ with 0s added to the bottom, up to the second-to-last element (which for $b \geq 3$ is $1_{\{2,b\}}$). Then run the sequence backward, except with 1s added to the bottom, and swapping coordinates 2 and $a - 1$ (the red and blue rows in the figure). Finally, finish the sequence from $1_{\{1,a\}}$ by walking through $1_{\{3,a\}}, 1_{\{4,a\}}, \dots, 1_{\{a-2,a\}}, 1_{\{2,a\}}$, with the appropriate weight-three sequences $1_{\{1,3,a\}}, 1_{\{3,4,a\}}, \dots, 1_{\{a-1,a-2,a\}}, 1_{\{a-2,2,a\}}$ (shown in gray) interposed.

To ensure that no vertex is visited more than once, one need only check that the last $2a - 5$ sequences are distinct from those that came before. For this, one can track by induction the $2a - 3$ hypercube vertices that are not visited by each walk: 0^a , the $a - 2$ weight-one strings $1_2, \dots, 1_{a-1}$, and the $a - 2$ weight-two strings $1_{\{1,3\}}, 1_{\{3,4\}}, 1_{\{4,5\}}, \dots, 1_{\{a-1,a\}}$. Thus, the sequence has total length $2^a - (2a - 3)$. $\square$

The length $2^a - 2a + 3$ is maximal. This follows since there are $2^{a-1} - a$ vertices with odd weight more than one, and vertices must alternate odd and even weights.

3 Distance-three stabilizer measurement

In this section, we outline two protocols for distance-three CSS fault-tolerant stabilizer measurement. They differ based on the speed of qubit measurement and reset.

For $w \in \{4, 5, 6\}$, flag-fault-tolerant circuits are constructed the same way regardless of qubit reset speed. We show in Fig. 4(a) that for $w = 6$, only two flag qubits are required. Lower-weight stabilizers can be measured by removing data CNOTs and making appropriate changes to the Pauli corrections. For $7 \leq w \leq 10$, the different constructions yield the same circuits. It is only for $w > 10$ that the effects of qubit reset speed are pronounced.

3.1 Fast reset

Theorem 3. If qubits can be measured and reset quickly, then for any w , four ancilla qubits are sufficient to measure the syndrome of $X^{\otimes w}$, CSS fault-tolerantly to distance three. Moreover, the number of measurements needed is $\lceil \frac{w+2}{4} \rceil + 1$.

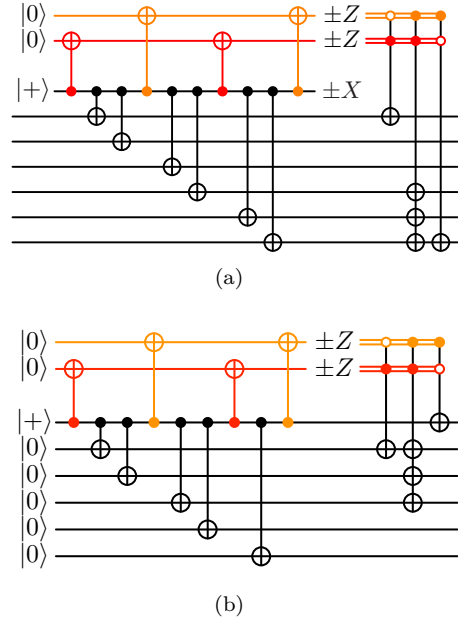


Figure 4: (a) Circuit to measure an $X^{\otimes 6}$ stabilizer, CSS fault-tolerant to distance three. (b) Circuit to prepare a six-qubit cat state, fault-tolerant to distance three.

Proof. For $w \in \{4, 5, 6\}$, the circuit using two flag ancillas is shown in Fig. 4(a). It runs through a sequence of three flag patterns and a multi-qubit correction is only applied for the flag pattern 11. For $w > 6$, the general construction is shown in Fig. 5. Each repetition of the highlighted region adds the X parity of four more data qubits, while measuring and quickly reinitializing one flag qubit. In terms of the number of measurements m , the construction achieves up to $w = 4(m - 1) - 2$. It is fault-tolerant because X faults on the control wire cause flag patterns of alternating weights two or three, that localize the fault to three possible consecutive locations along the control wire: before, between or after two CNOT gates. The appropriate correction, ensuring distance-three fault tolerance, is for a fault between the CNOT gates. $\square$

Theorem 3 may be optimal; it does not appear to be possible to use fewer than three flag qubits. With just one flag qubit, one can detect that an error has occurred, but not where. As illustrated in Fig. 6, either the control wire is unprotected at some point or for $w \geq 4$ there is no consistent correction rule.

By a similar argument, two flag qubits are not enough. Any correction based on a single flag can have weight at most one, since the flag measurement itself could be faulty. However, if at some point in the middle the control wire is protected by just a single flag, a weight-one correction will not suffice. On the other hand, if both flags are used to protect the control wire across the entire sequence of CNOT gates, we are unable to locate faults well enough to correct them.

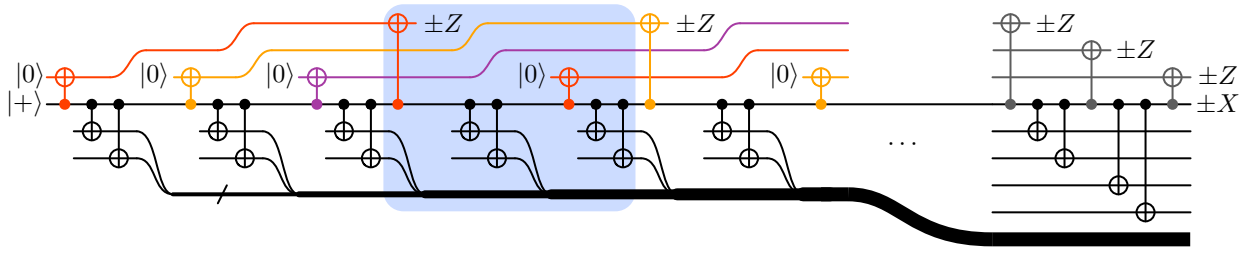


Figure 5: Distance-three fault-tolerant syndrome bit measurement only needs three flag qubits. The highlighted region can be repeated to fit the weight of the stabilizer being measured.

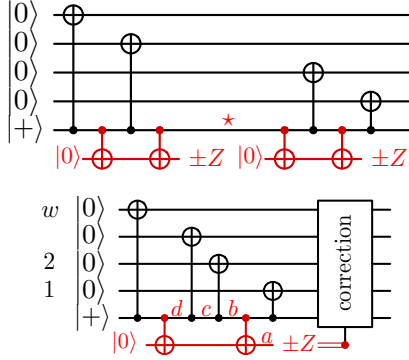


Figure 6: Distance-three error correction is not possible with one flag qubit. Either (top) the control wire is unprotected at some point $\star$, from which an X fault can propagate to an error of weight at least two; or (bottom) faults at a, b, c, d , causing respective errors I, X_1, X_1X_2, X_w have no consistent correction.

We remark that this construction can also be used to prepare a w -qubit cat state fault-tolerantly to distance three. The conversion follows three steps: 1. Remove one data qubit. 2. Initialize the data qubits as $|0\rangle$. 3. Remove the syndrome ancilla measurement, so as to retain it in the support of the stabilizer. An example of this conversion is shown for $w = 6$ in Fig. 4(b). In Section 4, we will give a better protocol that uses just one ancilla qubit.

3.2 Slow reset

Theorem 4. *The syndrome of $X^{\otimes w}$ can be measured CSS fault-tolerantly to distance three using $m \geq 3$ measurements, provided that*

$$w \leq 2(2^{m-1} - 2(m-1) + 3).$$

Proof. Two examples are shown in Fig. 4(a), for $w = 6$, and Fig. 1(b), for $w = 10$. As in these figures, in general we collect the syndrome two qubits at a time into a syndrome qubit that is initialized as $|+\rangle$. Between each of these pairs of CNOT gates, a CNOT is applied from the syndrome qubit into one of $m-1$ flag qubits. This leads to a sequence of flag patterns, e.g., 100, 110, 111, 011, 001 for the $w = 10$ example. Based on the observed flag pattern, a correction is

applied as if an X fault had occurred between the corresponding pair of flag CNOT gates.

Observe that the flag sequence changes one bit at a time; it can be thought of as a path on the hypercube. It begins and ends with weight-one patterns, but otherwise the patterns all have weight at least two. This is important for distance-three fault tolerance because a fault could affect the flags, and only the first and last data corrections have weight one. Also, the flag patterns along the sequence are distinct, so each is associated with only one correction. The theorem then just follows using the flag sequence construction in Lemma 2. $\square$

Note that the approach of Theorem 4, with slow reset, is different from the fast reset case of Theorem 3, in that a flag qubit is active and able to detect faults in more than one region of the circuit.

4 Distance-three cat state preparation

Next we turn to the question of distance-three fault-tolerant preparation of cat states. For preparing a two- or three-qubit cat state, any preparation circuit is automatically fault-tolerant, because every error has weight zero or one. For example, on three qubits $XXI \sim IIX$, since XXX is a stabilizer. Fault tolerance becomes interesting for preparing cat states on $w \geq 4$ qubits.

The ideas of Theorems 3 and 4 can also be applied to cat state preparation. For example, just as in Fig. 4 a circuit for measuring $X^{\otimes 6}$ with three ancilla qubits corresponds to a circuit to prepare a six-qubit cat state with two ancillas, similarly adapting the construction of Theorem 4 allows preparing a $2(2^a - 2a + 3)$ -qubit cat state using a ancilla qubits each measured once. However, we can do better.

Theorem 5. *For $m \geq 2$, one ancilla qubit, measured m times, is sufficient to prepare a cat state on w qubits fault-tolerantly to distance three, for*

$$w \leq 3(2^m - 2m + 2).$$

Let $[m] = \{1, 2, \dots, m\}$ and $X_S = \prod_{j \in S} X_j$.

Proof of Theorem 5. Figure 7 illustrates our construction for the cases $m = 3$ and $m = 4$. In general, we prepare a w -qubit cat state using CNOT gates from the first qubit, so that the possible X errors from a single fault are $I, X_1, X_{[2]}, X_{[3]}, \dots$. We then compute parities of subsets of the qubits into the ancillas, following the flag sequence from Lemma 2 and Fig. 3. Although for clarity Fig. 7 shows the m parity checks being made in parallel, they can also be made sequentially with just one ancilla qubit.

With the given correction rules, errors due to single faults are corrected up to possibly a weight-one remainder. (For example, in Fig. 7(a), errors $X_{[5]}, X_{[6]}$ and $X_{[7]}$ all result in the parity checks 111, for which the correction $X_{[6]}$ is applied.) The circuit also tolerates faults within the parity-check sub-circuit, because a single fault here can flip at most one parity, and no correction is applied for the weight-one patterns. $\square$

By this method, the cat state is prepared in depth $w - 1$. The depth of the parity check circuit increases exponentially as 2^{m-2} for $m \geq 3$ if we consider slow reset ($a = m$). This is evident from the flag sequences in Fig. 3 as the maximum number of times any flag bit is switched. The total depth of the circuit is then $(w - 1) + 2^{m-2}$.

Note that the construction from Theorem 5 does not help for syndrome measurement, because the parity checks would in general become entangled with the data.

We can do slightly better if we allow an *adaptive* circuit, in which the parity checks are chosen based on the outcome of a flag qubit measurement. For example, Figure 8 gives a circuit to prepare a 15-qubit cat state using $m = 3$ measurements. Here, the result of measuring the red ancilla determines how the other two ancillas are used.

Theorem 6. *Using an adaptive circuit, for $m \geq 2$, one ancilla qubit, measured m times, can be used to prepare a cat state on w qubits fault-tolerantly to distance three, for*

$$w \leq 3(2^m - 2m + 3).$$

Proof. Our construction will follow the same basic structure as the circuit in Fig. 8. Prepare the w data qubits as $|+0^{w-1}\rangle$, then apply $\text{CNOT}_{1,w}, \text{CNOT}_{1,w-1}, \dots, \text{CNOT}_{1,2}$ to get a cat state. Let $k = 3(2^{m-1}) - 2$. Just before $\text{CNOT}_{1,k+1}$ and just after $\text{CNOT}_{1,2}$, apply CNOTs into the first ancilla qubit, the red qubit in Fig. 8, and measure it.

The remainder of the circuit depends on the measurement result. If it is 1, then a fault has been detected. The error on the cat state can be one of

$$I, X_1, X_{[2]}, X_{[3]}, X_{[4]}, X_{[5]}, \dots, X_{[k-1]}, X_{[k]}, X_{[k+1]}.$$

The correction procedure needs to determine in which of the above $1 + \frac{k-1}{3}$ groups-of-three the error lies;

then for any error in $\{X_{[3j]}, X_{[3j+1]}, X_{[3j+2]}\}$ the correction $X_{[3j+1]}$ works. Perhaps the easiest way to locate the error is by binary search using the Gray code in Lemma 1, e.g., by computing parities between qubits $3j$ for $j \in \{1, 2, \dots, 1 + \frac{k-1}{3}\}$. Since the measurement of the red ancilla could have been incorrect, it is important that the all-0s outcome of the binary search correspond to the $I, X_1, X_{[2]}$ error triple, as in Table 2. Using $m - 1$ measurements, we can search 2^{m-1} possibilities, which indeed is $1 + \frac{k-1}{3}$. (The search circuit can also be made nonadaptive, as in Fig. 8.)

Next consider the case that the first measurement result is 0, so no fault has been detected. The error on the cat state can be one of $X_{[k+1]}, X_{[k+2]}, \dots, X_{[w]} \sim I$. We again use the remaining $m - 1$ ancilla qubits to measure parities of subsets of cat state qubits. Since there is no guarantee of a fault having occurred yet, we use flag sequences from Lemma 2, where the length of the weight-at-least-two flag sequence is $J = 2^{m-1} - 2(m - 1) + 1$. The parity checks are now done between qubits $\{k, k + 1 + 3j, k + 2 + 3j\}$ for $j \in \{0, 1, \dots, J\}$, as shown in Fig. 9 and Table 2. We do not allow weight-one flag patterns to be able to correct any errors since they may also be triggered by a measurement fault on any one of the data qubits involved in the parity check.

Consolidating, we are allowed up to $3J + 1$ CNOTs before the red ancilla is initialized, and up to k CNOTs in the monitored region of the red ancilla. In total we can create a cat state on up to

$$w \leq 3J + k + 2 = 3(2^m - 2m + 3)$$

qubits, with m total measurements. $\square$

We also tested protocols where multiple flags are used for the initial partial localization of a fault (in place of the red flag qubit). We found no improvement to our bounds on ancilla overhead. It appears that ancillas are better used in the parity checks than for partial fault localization.

5 Simulation and space-time cost

We count the circuit depth and number of ancillas used in our distance-three fault-tolerant stabilizer measurement circuits. Parallelization can substantially reduce circuit depth. Table 3 compares our flag method for measuring a weight- w stabilizer to the earlier methods in Fig. 2. Also considered is a parallelized Shor method, in which the initial cat state is prepared in logarithmic depth, with $w/4$ extra ancilla qubits for postselection checks. The Shor methods must pass the postselection checks, and so they are non-deterministic protocols. Table 3 shows the best case, where all the checks pass. Note that the flag and parallelized Shor methods both have space $\times$ depth cost scaling as $O(w \log w)$, with the leading coefficient in favor of the flag method.

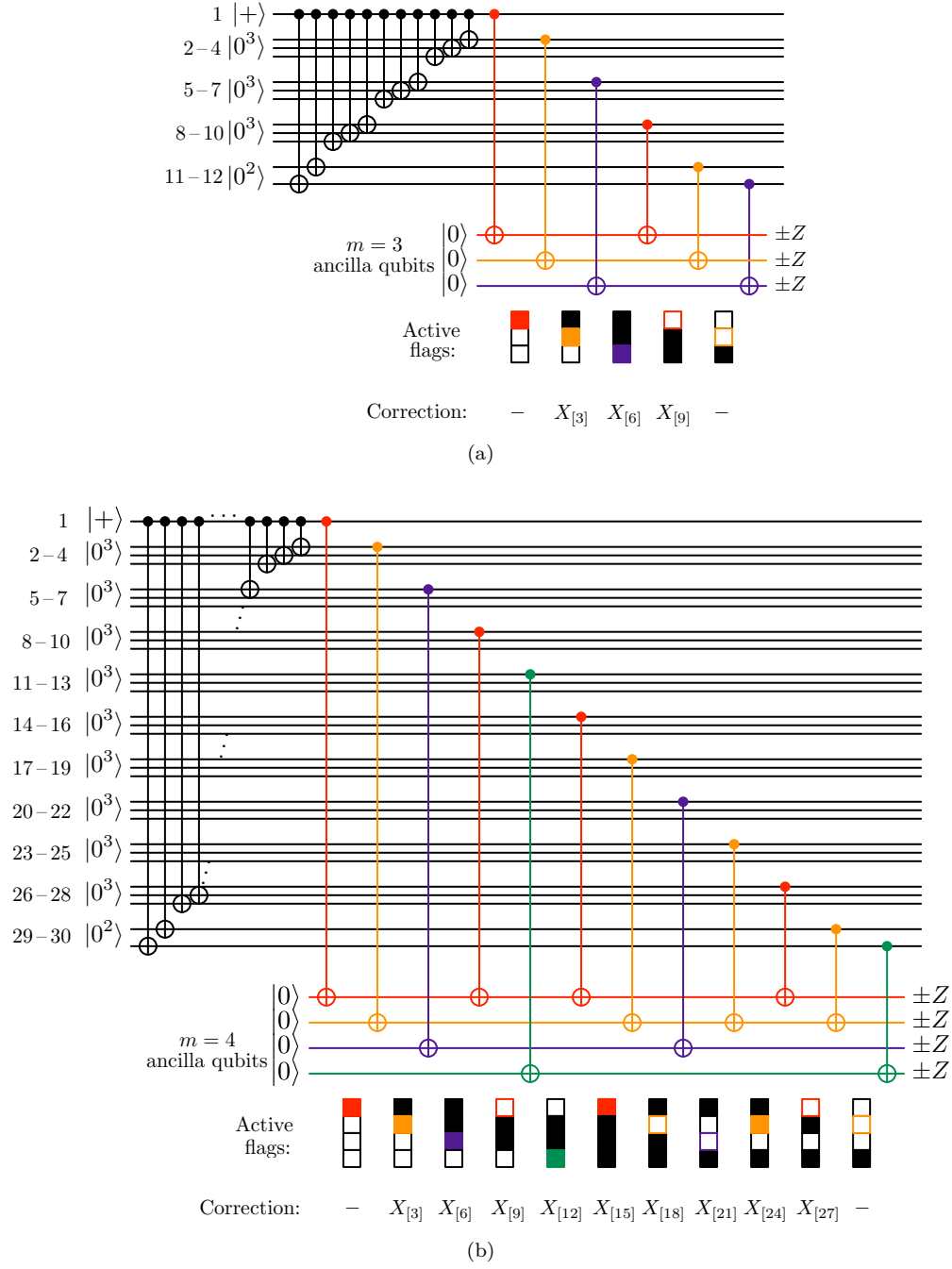


Figure 7: Distance-three fault-tolerant cat state preparation circuits. Note that, with fast reset, only one ancilla qubit is required.

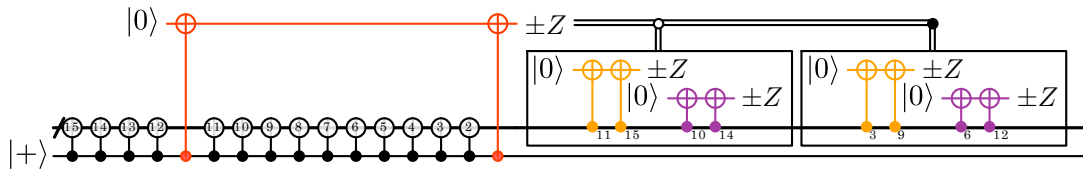


Figure 8: Circuit to prepare a 15-qubit cat state by adaptive error correction, fault-tolerant to distance three. Labels on the thick black wire indicate which data qubit in the block is being addressed as the control or target of the CNOT. If a fault occurs while preparing the cat state on the $|+\rangle$ qubit, it is partially localized by the red flag ancilla. The measurement result of this flag then determines a set of parity checks to completely localize a possible fault. After all the ancilla qubits have been measured, corrections are applied based on Table 2.

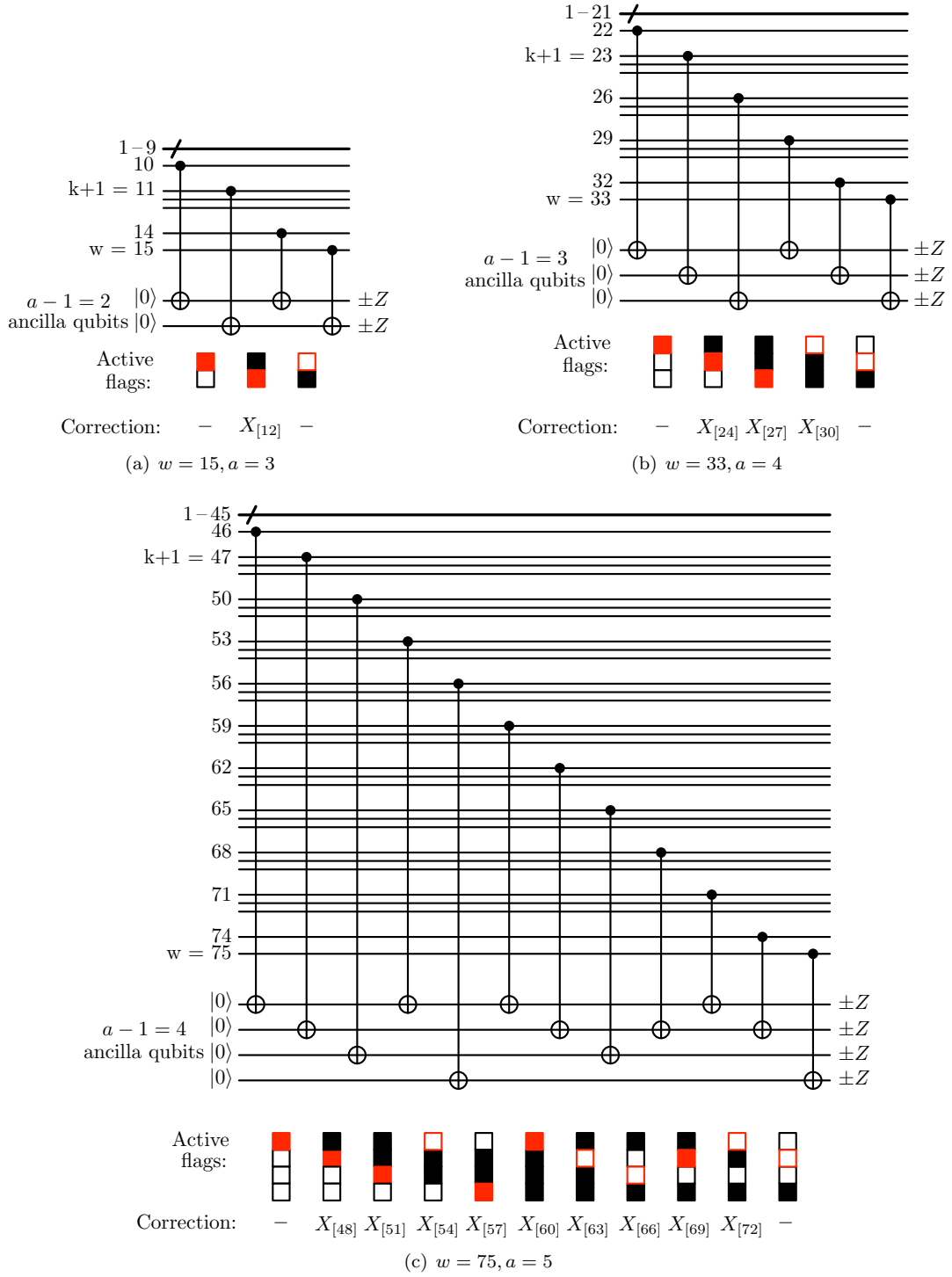


Figure 9: If the red ancilla flag in Fig. 8 is not triggered, these circuits are used to find and correct a possible error. The flag sequences (from Fig. 3) and corresponding corrections are listed at the bottom. Note that these sequences are nonadaptive, and can be used either with a ancilla qubits in a slow reset model, or with just one ancilla qubit in a fast reset model, since all the CNOT gates commute.

Table 2: Possible data errors and associated corrections for the different observed flag patterns in Fig. 8. $[m] = \{1, 2, \dots, m\}$.

Red flag	Parity checks		Possible errors	Correction
1	$3 \oplus 9$	$6 \oplus 12$		
	0	0	$I, X_1, X_{[2]}$	X_1
	1	0	$X_{[3]}, X_{[4]}, X_{[5]}$	$X_{[4]}$
	1	1	$X_{[6]}, X_{[7]}, X_{[8]}$	$X_{[7]}$
	0	1	$X_{[9]}, X_{[10]}, X_{[11]}$	$X_{[10]}$
0	$11 \oplus 15$	$10 \oplus 14$		
	0	0	I	None
	1	0	$X_{11}, X_{15}, X_{[14]}$	None
	1	1	$X_{[11]}, X_{[12]}, X_{[13]}$	$X_{[12]}$
	0	1	X_{10}, X_{14}	None

Table 3: Space and time costs for measuring a weight- w stabilizer using different distance-three fault-tolerant stabilizer measurement circuits. In the following, all the logarithms are base 2. The flag method requires the fewest ancillas and has low depth, allowing for the smallest cost when computing $\# \text{ancillas} \times \text{depth}$.

Protocol	Ancillas	Depth	Ancillas $\times$ Depth
Shor	$w + 1$	$w/2 + 3$	$O(0.5w^2)$
Shor-Par	$5w/4$	$3 \log w - 1$	$O(3.75w \log w)$
DA	w	$2w - 1$	$O(2w^2)$
Compressed DA	$w/2$	$3w/2 - 2$	$O(0.75w^2)$
Flag	$\log w + 1$	$3w/2 + O(1)$	$O(1.5w \log w)$
Not fault-tolerant	1	w	$O(w)$

Using a standard depolarizing noise model, we simulate noisy versions of the different circuits to determine statistics of the weight-one and weight-two errors. Specifically:

- With probability p , the preparation of $|0\rangle$ is replaced by $|1\rangle$ and vice versa—similarly for $|+\rangle$ and $|-\rangle$.
- With probability p , an X or Z measurement has its outcome flipped.
- With probability p , a one-qubit gate is followed by a Pauli error drawn uniformly at random from $\{X, Y, Z\}$.
- With probability p , the two-qubit CNOT gate is followed by a two-qubit Pauli error drawn uniformly at random from $\{I, X, Y, Z\}^{\otimes 2} \setminus \{I \otimes I\}$.

There are no errors on idle resting qubits.

Figure 10 shows the rates of weight-one errors and weight-two errors for different input physical error rates p . The rate of weight-one errors is lowest in the non-fault-tolerant circuit, since it contains the fewest locations for faults. The Shor method has a lower weight-one error rate than the flag method, but

among the deterministic fault-tolerant methods, the flag method performs the best. For larger stabilizers ($w = 22$), the curves for the Shor method and the flag method are closer, implying that the difference in the rate of weight-one errors between the Shor method and the flag method is reduced.

As expected, the rate of weight-two errors of the three fault-tolerant protocols scales quadratically with p , allowing for a lower probability of weight-two errors below a pseudothreshold (physical error rate below which a fault-tolerant method achieves lower weight-two error rate than the non-fault-tolerant method). Notice that the flag method has the highest pseudothreshold. Moreover, as the stabilizer weight is increased, the pseudothreshold of the flag method decreases slower than those of the other fault-tolerant methods. Asymptotically, the flag method admits the highest pseudothreshold for weight-two errors, but incurs more weight-one errors than the probabilistic Shor method. Additionally, we compute the rate of errors on the syndrome bit, as this determines how much fault tolerance will be needed to correct faulty syndrome information [17]. The rate of faulty syndrome bits is lowest when using the flag method for fault tolerance.

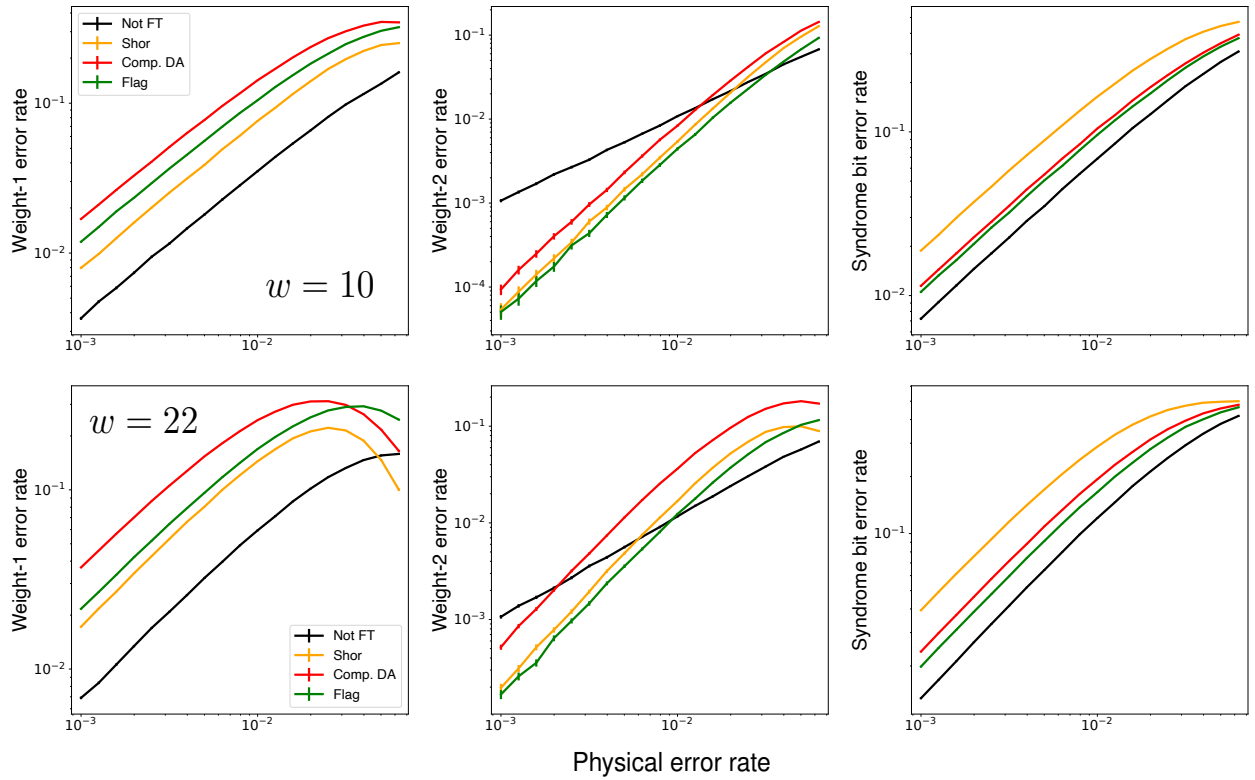


Figure 10: Simulation of the noisy measurement of an $X^{\otimes 10}$ and $X^{\otimes 22}$ stabilizer at physical error rate $p \in \{10^{-3}, 10^{-2}\}$ using different distance-three fault-tolerant circuits: Shor-style, compressed Divincenzo-Aliferis, and the flag method of Section 3.2. In the first and second column of graphs, we show the rate of weight-one and weight-two data errors due to these circuits, with 99% error bars. In the third column, we show the rate at which the measured syndrome bit is wrong.

6 Conclusion

In this paper, we optimize the overhead of distance-three fault-tolerant stabilizer measurement and cat state preparation. If the circuit on w qubits must tolerate one fault, we show that $\sim \log w$ extra qubits are sufficient. We detail the construction of a maximal-length path on the hypercube and show that, compared to previous flag schemes, it allows using the extra flag qubits more efficiently to catch and distinguish faults.

We describe two circuits for stabilizer measurement based on the speed of ancilla qubit reset. With slow reset, a weight- w stabilizer can be measured fault-tolerantly to distance three using only $\lceil \log_2 w \rceil$ flag qubits for fault tolerance. With fast reset, only three flag qubits are required, but the number of times they are measured and reset grows as $w/4$.

In our circuits for fault-tolerant cat state preparation we check for errors after the cat state is non-fault-tolerantly prepared. We show, using a deterministic and an adaptive circuit, that the overhead for fault tolerance can be as low as logarithmic in the size of the cat state. In fact, only one flag qubit suffices, as long as it can reset quickly.

We now turn to further improvements. The circuits detailed in this paper are only fault-tolerant to distance three. However, flags can be used to effect

fault tolerance to arbitrary distance [18, 19], and it is open to develop higher-distance fault-tolerant stabilizer measurement circuits with low overhead.

From the perspective of stabilizer algebra, a cat state is a CSS ancilla state. A future avenue of research is to extend these flag techniques to fault-tolerantly and deterministically prepare more complex CSS ancilla states.

In order to execute the circuits in this paper, one qubit needs to be connected to all the other qubits used. This is concerning for architectures with limited connectivity, such as superconducting qubits. But by mixing flag and transversal gate concepts for fault tolerance, it is possible to construct stabilizer measurement circuits that can measure arbitrarily large stabilizers using only local interactions, fault-tolerantly.

7 Acknowledgements

We would like to thank Rui Chao, Sourav Kundu and Zhang Jiang for insightful conversations. Research supported by Google and by MURI Grant FA9550-18-1-0161. This material is based on work supported by the U.S. Department of Energy, Office of Science, National Quantum Information Science Research Centers, Quantum Systems Accelerator.

References

- [1] Daniel M. Greenberger, Michael A. Horne, and Anton Zeilinger. “Going beyond bell’s theorem”. Pages 69–72. Springer Netherlands. Dordrecht (1989). [arXiv:0712.0921](#).
- [2] Nicolas Gisin and Rob Thew. “Quantum communication”. *Nature Photonics* **1**, 165 (2007). [arXiv:quant-ph/0703255](#).
- [3] Jian-Wei Pan, Zeng-Bing Chen, Chao-Yang Lu, Harald Weinfurter, Anton Zeilinger, and Marek Żukowski. “Multiphoton entanglement and interferometry”. *Rev. Mod. Phys.* **84**, 777 (2012). [arXiv:0805.2853](#).
- [4] P. W. Shor. “Fault-tolerant quantum computation”. In Proceedings of the 37th Annual Symposium on Foundations of Computer Science. Page 56. FOCS ’96USA (1996). IEEE Computer Society. [arXiv:quant-ph/9605011](#).
- [5] David P. DiVincenzo and Panos Aliferis. “Effective fault-tolerant quantum computation with slow measurements”. *Phys. Rev. Lett.* **98**, 020501 (2007). [arXiv:quant-ph/0607047](#).
- [6] Ashley M. Stephens. “Efficient fault-tolerant decoding of topological color codes” (2014). [arXiv:1402.3037](#).
- [7] Theodore J. Yoder and Isaac H. Kim. “The surface code with a twist”. *Quantum* **1**, 2 (2017). [arXiv:1612.04795](#).
- [8] Christopher Chamberland, Aleksander Kubica, Theodore J Yoder, and Guanyu Zhu. “Triangular color codes on trivalent graphs with flag qubits”. *New Journal of Physics* **22**, 023019 (2020). [arXiv:1911.00355](#).
- [9] Christopher Chamberland, Guanyu Zhu, Theodore J. Yoder, Jared B. Hertzberg, and Andrew W. Cross. “Topological and subsystem codes on low-degree graphs with flag qubits”. *Phys. Rev. X* **10**, 011022 (2020). [arXiv:1907.09528](#).
- [10] Rui Chao and Ben W. Reichardt. “Quantum error correction with only two extra qubits”. *Phys. Rev. Lett.* **121**, 050502 (2018). [arXiv:1705.02329](#).
- [11] Rui Chao and Ben W. Reichardt. “Fault-tolerant quantum computation with few qubits”. *npj Quantum Information* **4**, 42 (2018). [arXiv:1705.05365](#).
- [12] A. M. Steane. “Active stabilization, quantum computation, and quantum state synthesis”. *Phys. Rev. Lett.* **78**, 2252 (1997). [arXiv:quant-ph/9611027](#).
- [13] E. Knill. “Scalable quantum computing in the presence of large detected-error rates”. *Phys. Rev. A* **71**, 042322 (2005). [arXiv:quant-ph/0312190](#).
- [14] Shilin Huang and Kenneth R. Brown. “Between shor and steane: A unifying construction for measuring error syndromes”. *Phys. Rev. Lett.* **127**, 090505 (2021).
- [15] Frank Gray. Pulse code communication. U.S. Patent 2632058A. Issued Mar. 17, 1953.
- [16] M. Gardner. “The Binary Gray Code”. In *Knotted Doughnuts and other Mathematical Entertainments*. Pages 22–39. W. H. Freeman and Company, New York (1986).
- [17] Nicolas Delfosse, Ben W. Reichardt, and Krysta M. Svore. “Beyond single-shot fault-tolerant quantum error correction”. *IEEE Transactions on Information Theory* **68**, 287–301 (2022). [arXiv:2002.05180](#).
- [18] Christopher Chamberland and Michael E. Beverland. “Flag fault-tolerant error correction with arbitrary distance codes”. *Quantum* **2**, 53 (2018). [arXiv:1708.02246](#).
- [19] Rui Chao and Ben W. Reichardt. “Flag fault-tolerant error correction for any stabilizer code”. *PRX Quantum* **1**, 010302 (2020). [arXiv:1912.09549](#).

A Postselected cat state preparation tolerating two faults

Shor’s method for fault-tolerant stabilizer measurement relies on the fault-tolerant preparation of a cat state by postselection. In Fig. 2(a), the cat state is prepared fault-tolerantly to distance-two; it detects one fault. For postselected distance-three fault tolerance, any one or two faults in the circuit must result in an error of weight at most one or two respectively, else the state must be rejected. In Fig. 11 we show how to prepare a weight-12 cat state fault-tolerantly to distance three—detecting up to two faults.

Theorem 7. *One ancilla qubit measured $m \geq 2$ times, can be used to prepare a cat state on w qubits fault-tolerantly to distance three, detecting up to two faults, for*

$$w \leq 3 \cdot 2^{m-1}.$$

Proof. We explain the proof using the circuit in Fig. 11. The circuit passes with acceptable weight-one or weight-two errors when all the flag qubits are measured as 0. If one X fault occurs on the $|+\rangle$ qubit during the preparation of the cat state, it may spread to a data error of weight > 1 . However the red flag qubit is triggered and the fault is detected. If two X faults occur on the $|+\rangle$ qubit, the red flag qubit may not catch it, yet a data error of weight > 2 can exist on the cat state. Since this scenario only arises from two faults, it suffices to check the parities between every third qubit of the cat state, as an error on two

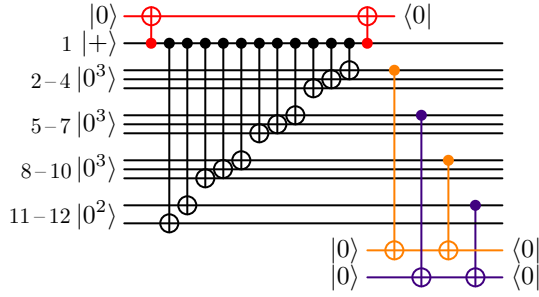


Figure 11: Two-error-detecting fault-tolerant circuit for the preparation of a weight-12 cat state. The state is only accepted when all flag qubits are measured as 0. Note that with fast reset, only one ancilla qubit is required.

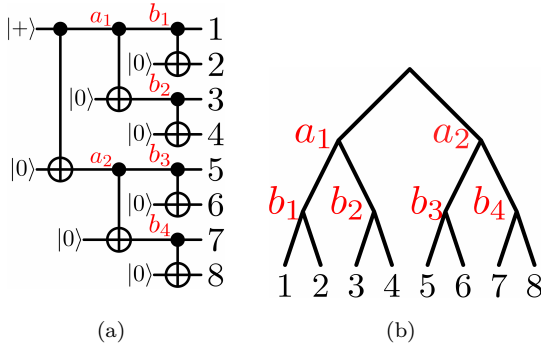


Figure 12: (a) Logarithmic-depth preparation of an eight-qubit cat state shows there are six possible locations for X faults that create errors of weight at least two. Parity checks need to be chosen to find corrections that leave the cat state with error of weight less than two. (b) The circuit on the left can be represented as a graph, where a CNOT gate is represented by the splitting of an edge.

consecutive qubits is acceptable. Higher-weight errors, such as the weight-seven error $X_2 X_3 \dots X_8$ in Fig. 11 may not be detected by parity checks that have an even number of erroneous qubits. However these errors are always caught by other parity checks.

To check for errors of weight greater than two, we perform parity checks similar to that in Theorem 5. Instead of the flag sequence from Lemma 2, the Gray code from Lemma 1 is used. Now the parities are computed between qubits $3j - 1$ for $j \in \{1, 2, \dots, 2^{m-1}\}$. The first and the last qubits are not checked for errors and so with m flags, the maximum cat state weight achieved is $3 \cdot 2^{m-1}$. $\square$

B Parallelized distance-three cat state preparation

So far, we have focused on fault-tolerant preparation circuits with depth linear in the cat state weight. In this section, we detail how to prepare cat states fault-tolerantly in logarithmic depth.

In Fig. 12(a) an eight-qubit cat state is prepared in three rounds of CNOT gates. There are six locations (marked in red) where an X fault may cause an error of weight at least two. These faults result in data errors with a different structure from the linear-depth protocols of Section 4, hence different parity checks are required. It is simpler to determine these parity checks if the circuit is viewed as a binary tree, as in Fig. 12(b). Here time flows down and every CNOT onto a fresh $|0\rangle$ qubit is denoted by the splitting of an edge. An X fault at a marked location results in an X error on all the leaf nodes directly under the location. Note that a fault at the root cannot cause a bad error.

We use only two-qubit parity checks, however larger parity checks may be used at the expense of increased depth. If a parity check checks qubit x , it provides information on whether a fault occurred anywhere in the lineage: $l(x) = \{x, \text{parent}(x), \text{parent}(\text{parent}(x)), \dots, \text{root}\}$. Therefore, if a parity check (x, y) is triggered, a fault at one of the locations $l(x) \cup l(y) \cup \{\text{SPAM}\}$ has occurred, where $\{\text{SPAM}\}$ is the set of faults during state preparation or measurement of the parity-check qubit.

Using the parity checks $(1, 5), (2, 7), (3, 6), (4, 8)$, it is possible to separate the five distinct weight at least two errors (since the error due to a_1 and a_2 is the same up to the cat state's $X^{\otimes w}$ stabilizer) into distinct triggered flag patterns:

	(1, 5)	(2, 7)	(3, 6)	(4, 8)
a_1, a_2	•	•	•	•
b_1	•	•	○	○
b_2	○	○	•	•
b_3	•	○	•	○
b_4	○	•	○	•

Note that a fault at any of the above locations requires a multi-qubit data correction. We ensure that each of them is detected by at least two parity checks, as one faulty parity check must not induce corrections of weight greater than one.

Theorem 8. *Using parallelized circuits, a w -qubit cat state can be prepared fault-tolerantly to distance three using $\frac{w}{2}$ parity checks, where $\frac{w}{2} = 2^j$, $j \in \mathbb{N}$. The depth of the circuit is $2 + \log_2 w$.*

Proof. For parity check $i \in \{1, 2, \dots, \frac{w}{4}\}$, the cat state qubits checked are $(i, \frac{w}{2} + 2i - 1)$. For the remaining parity checks $i \in \{\frac{w}{4} + 1, \frac{w}{4} + 2, \dots, \frac{w}{2}\}$, the qubits checked are $(i, 2i)$. As in Fig. 12(b), faults at the a level (depth-one) locations trigger all the parity checks, since each parity check is executed on one cat state qubit from the first half, and one from the second. The correction $X^{\otimes w/2}$ on either half of the qubits works for both faults as $(X^{\otimes w/2} \otimes I^{\otimes w/2})(I^{\otimes w/2} \otimes X^{\otimes w/2}) = X^{\otimes w}$ is a stabilizer of the cat state. Faults at the b level (depth-two) trigger distinct sets of $\frac{w}{2}$ parity checks, where the correction is on all the leaf nodes

under the uniquely identified fault. The same holds for faults at depth- k , which trigger distinct sets of $\frac{w}{2^k}$ parity checks.

One faulty parity check leads to a weight-one flag pattern, for which we do not apply corrections, as the error is restricted to at most one cat state qubit. $\square$

C Distance-five and distance-seven fault-tolerant stabilizer measurement

Distance-five fault tolerance is interesting for stabilizers of weight $w \geq 6$. For $w \in \{6, 7, 8\}$, the circuits in Fig. 13 with seven ancilla qubits are distance-five fault-tolerant. We present a general method to construct stabilizer measurement circuits for arbitrary w in Fig. 14. By computer simulation, we verify the fault tolerance of this construction for w up to 90 qubits. The general construction proceeds as follows. First, five flag qubits are activated. For each additional flag that is needed, the gates in the shaded blue region are applied. These gates deactivate an existing flag and activate a new flag. Finally, when no additional flags are needed, the flags are deactivated in the order $\{2, 4, 1, 5, 3\}$. 1 denotes the flag that has been active for the longest time and 5, the flag that has been active for the shortest time. To ensure faults are correctly flagged, it is necessary to ensure there is asymmetry between the order in which flags are activated and the order in which they are deactivated. This is in contrast to the distance-three DiVincenzo-Aliferis method in Fig. 2(c), where both the orders are symmetric.

In Fig. 14, the thick black line indicates the w -qubit register of data qubits that are in the support of the stabilizer. Data CNOT gates (in black) are applied to qubits $\{w, w-1, \dots, 1\}$ after every flag CNOT (in red). The last data CNOT must be placed either before the third-last or second-last flag CNOT. The addition of another data CNOT gate before the last flag CNOT results in uncorrectable errors.

If there are a ancilla qubits, one can measure a weight- $(2a-5)$ or weight- $(2a-4)$ stabilizer. Hence a weight- w CSS stabilizer may be fault-tolerantly measured to distance-five, for $w \leq 2a-4$. Note that at most five flag qubits are active at any instant. Hence with fast qubit reset, one only requires five flag ancillas and one syndrome ancilla to measure an arbitrary weight stabilizer fault-tolerantly to distance-five.

For distance-seven fault-tolerance, we detail changes to the spacing between data CNOT gates and generalize the order in which flag ancillas are activated and deactivated. We show how to construct circuits for stabilizer of arbitrary weight w by first discussing a circuit for a weight-17 stabilizer, shown in Fig. 15. We chose $w = 17$ since the circuit is non-trivial and its construction encompasses all the tricks needed to construct circuits for arbitrary weight.

In general, compared to Fig. 14, the number of ancilla CNOT gates between data CNOT gates is doubled, except in the center of the circuit, where it is tripled for the length of four data CNOT gates. For odd w , the number of ancilla CNOTs between the $w-1$ subsequent pairs of data CNOT gates is the sequence $\{(\lceil \frac{w-6}{2} \rceil \text{ 2's}), 3, 3, 3, 3, (\lfloor \frac{w-6}{2} \rfloor \text{ 2's}), 1\}$, as shown in Fig. 15. For even w , the sequence is $\{(\frac{w-6}{2} \text{ 2's}), 3, 3, 3, 3, (\frac{w-6}{2} \text{ 2's}), 1\}$. Note that, as shown in Fig. 15, one additional ancilla CNOT gate is required at the start.

Next we comment on the order in which ancilla qubits are deactivated as flags. Similar to the distance-five case, after initially activating seven flags, a flag is deactivated to activate a new flag qubit. An active group of seven flags is closed in the order $\{2, 4, 6, 1, 3, 5, 7\}$. As these seven flags are closed, seven new flags are simultaneously opened. The process repeats until there are exactly seven remaining flags to close. These last seven flags are also closed in the same order $\{2, 4, 6, 1, 3, 5, 7\}$. In Fig. 15, flag ancillas are shown in alternating colors to highlight the order that flags are activated and deactivated. Distance-seven fault-tolerance was verified by computer simulation for stabilizer weight up to 32. The number of flag ancillas needed to measure a weight- w stabilizer is $w+1$.

The techniques described in this section may also be used to develop resource-efficient circuits that are fault-tolerant to higher distance.

