

Codes and designs for quantum error correction: some recent developments through combinatorics

Yuichiro Fujiwara *

Abstract

Suppressing the effect of decoherence plays a vital role in the theory of quantum information processing. Despite many skepticisms about its feasibility, the existence of error correction schemes in the quantum domain was proved in the last century, giving the birth to the field of quantum error correction. Since then, the new field has seen remarkably rapid progress in various directions including experimental realizations of quantum error-correcting codes. While quantum information science draws on many branches of physics, computer science, and mathematics, combinatorics is rapidly becoming an indispensable mathematical tool to the study of quantum error correction. In this note, we review some of the latest developments in quantum error correction where combinatorial design theory and coding theory play the central role.

1 Introduction

Quantum error-correcting codes are mathematical schemes for recovering the original quantum states of quantum bits, or *qubits*, that carry the information when the quantum states are transformed by unintended quantum operations, that is, *quantum noise* [40]. It is vital to suppress the effect of quantum noise for realizing large-scale quantum computation and quantum communication because qubits are highly vulnerable to environmental noise in practical and realistic situations.

While the importance of error correction is apparent in the quantum domain, there had been doubts about the existence of a feasible scheme in the quantum domain until the discovery of the famous 9-qubit code [43] and 7-qubit code [44] in the mid 1990's. Research on quantum error correction has seen rapid and remarkable progress since then. In fact, various types of quantum error-correcting codes are now known such as the celebrated stabilizer codes [20, 7], which encompasses the first two quantum error-correcting codes. Small quantum error-correcting codes have been experimentally realized as well [10, 32, 8, 5, 35, 2, 42, 37, 41, 51, 53, 52, 3].

However, the remarkable progress we have seen does not imply that the theory of quantum error correction has become as mature as that of classical error correction. For instance, while the stabilizer formalism developed in [21] has given rise to a wide range of quantum error-correcting codes, the admissible structures of a code are severely restricted compared to the freedom in designing classical error-correction codes. The fact that we do not have many successful general frameworks for quantum error correction also limits the variety of quantum error-correcting codes.

One effective way to overcome the limitations and difficulties is to develop a new framework for designing quantum error-correcting codes that makes it possible to directly import a wider range of classical coding theory to the quantum domain. A notable breakthrough in this direction is the *entanglement-assisted stabilizer formalism*, where one may fully exploit any binary or quaternary linear code over the binary field $\mathbb{F}_2$ or the finite field $\mathbb{F}_4$ of order four respectively for correcting errors on qubits, provided that there is an adequate supply of maximally entangled noiseless qubits, called *ebits* [6]. Entanglement-assisted quantum error-correcting codes can be regarded as generalized stabilizer codes in the sense that those codes requiring no ebit are the standard stabilizer codes. In other words, if a given linear code can not be turned into a quantum error-correcting code through the standard stabilizer formalism, it may still be fully exploitable if some ebits can be shared through a noiseless channel to help encode and decode noisy qubits.

*Division of Physics, Mathematics and Astronomy, California Institute of Technology, MC 253-37, Pasadena, California 91125, USA
(email: yuichiro.fujiwara@caltech.edu)

A downside of entanglement assistance is that completely noiseless qubits are extremely difficult to realize in a practical quantum device. This disadvantage is particularly pronounced in the context of protecting stored quantum information, where the information source and sink are separate in the time domain. This characteristic of entanglement-assisted quantum error-correcting codes led to a series of research that tried to identify excellent linear codes which can be imported by using only a small number of ebits [25, 47, 24, 12, 11, 15, 48, 49, 26, 34, 18, 23].

Very recently, a framework that significantly reduces this burden of providing noiseless qubits was proposed, where any binary or quaternary linear codes over $\mathbb{F}_2$ or $\mathbb{F}_4$ respectively can be fully exploited as long as one can provide auxiliary qubits that are only subject to a restricted quantum error model [14]. This framework exploits the fact that while realizing completely noiseless qubits is very difficult, not every kind of quantum error is equally difficult to suppress through technical advances on hardware. For example, it is known that one can correct any type of quantum error in the standard error model if two particular types of error, called a bit error and phase error, can be corrected under the assumption that both may happen on the same qubit [40]. However, phase errors are expected to be far more likely than bit errors in most actual quantum devices [28]. Under the newer framework, one may choose an error model where most qubits suffer from bit errors and phase errors while only phase errors may occur on the auxiliary qubits. Hence, unlike the entanglement-assisted stabilizer formalism, which requires completely noiseless qubits, the newer framework only needs more easily achievable “less noisy” ones.

This note gives a brief summary of the latest framework and presents a short account on how coding theory and design theory contribute to quantum information science through this recent idea.

2 Quantum Error Correction with Reliable Auxiliary Qubits

Here we give a brief review of how reliable auxiliary qubits help correct quantum errors in the recent frameworks. For the basics of quantum information theory and classical coding theory, we refer the reader to [40, 27].

As usual, by a binary linear $[n, k, d]$ code, we mean a k -dimensional subspace $\mathcal{C}$ of the n -dimensional vector space over $\mathbb{F}_2$ in which a nonzero vector with the smallest number of nonzero entries has exactly d nonzero entries, that is, $\min\{\text{wt}(\mathbf{c}) \mid \mathbf{c} \in \mathcal{C}, \mathbf{c} \neq \mathbf{0}\} = d$. Because we only consider a binary code, we omit the term binary when referring to linear codes and LDPC codes. A linear $[n, k, d]$ code encodes k logical bits into n physical bits, where n is the *length* and k is the *dimension* of the code. The parameter d is the *minimum distance* of the code and ensures that up to $\lfloor \frac{d-1}{2} \rfloor$ errors can be corrected. Analogously, an $[[n, k, d]]$ quantum error-correcting code of *length* n , *dimension* k , and *distance* d encodes k -qubit information into n physical qubits and corrects up to errors on $\lfloor \frac{d-1}{2} \rfloor$ qubits.

A fundamental fact in the quantum domain is that, through a process called *discretization*, an error correction scheme can correct any general quantum error on one qubit if it can correct the effects of the Pauli operators X , Z and their product XZ , where the operator X corresponds to the *bit error* on one qubit while Z represents the *phase error* [40]. Similarly, quantum errors on multiple qubits are correctable if the corresponding transformation by a combination of X , Z and both at the same time on each of the affected qubits is correctable. The quantum error-correcting codes we consider in this note also exploit discretization. Hence, without loss of generality, we always assume that a quantum channel may introduce on each qubit only a bit error, a phase error or both at the same time as a quantum error.

The following is the fundamental tool we use to import linear codes.

Theorem 2.1 ([14]) *If there exists a linear $[n, k, d]$ code, then there exist unitary operations that encode k logical qubits into $2n - k$ physical qubits and correct up to $\lfloor \frac{d-1}{2} \rfloor$ quantum errors under the assumption that a fixed set of $2(n - k)$ physical qubits may experience phase errors but no bit errors.*

The point of the above theorem is that any linear $[n, k, d]$ code, which corrects errors on up to $\lfloor \frac{d-1}{2} \rfloor$ bits, can be turned into a $[[2n - k, k]]$ quantum error-correcting code that corrects quantum errors on up to $\lfloor \frac{d-1}{2} \rfloor$ qubits as long as predetermined $2(n - k)$ qubits are only subject to phase errors. Note that if the linear $[n, k, d]$ code is of sufficiently high rate so that k is close to n , the $2(n - k)$ auxiliary qubits consist of only a small fraction of the $2n - k$ physical qubits.

A useful fact regarding this type of quantum error correction is that we can employ decoding methods for linear codes based on error syndromes.

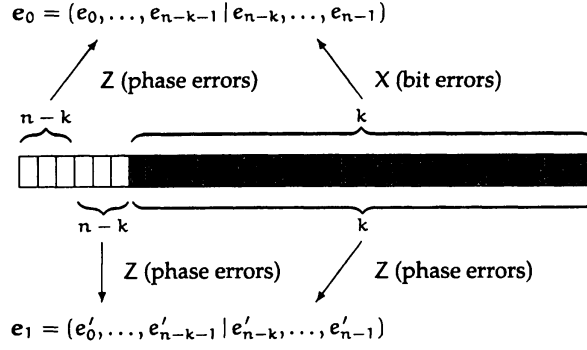


Figure 1: Correspondence of types and positions of quantum errors to error vectors. | The white boxes represent the $2(n-k)$ less noisy qubits that may experience only phase errors. The gray boxes are the k noisy qubits that may suffer from bit and/or phase errors. The first $n-k$ bits of e_0 and the first $n-k$ bits of e_1 correspond to whether phase errors occurred on the $2(n-k)$ less noisy qubits. The remaining k bits of e_0 indicate whether bit errors occurred on the k noisy qubits while the remaining k bits of e_1 correspond to phase errors on these noisy qubits.

Theorem 2.2 [17] Let $\mathcal{C}$ be a linear $[n, k, d]$ code. Assume that $2n-k$ physical qubits q_i , $0 \leq i \leq 2n-k-1$, are sent through a noisy quantum channel in which the first $2(n-k)$ qubits q_i , $0 \leq i \leq 2(n-k)-1$ are only subject to phase errors while the remaining k qubits q_i , $2(n-k) \leq i \leq 2n-k-1$ are subject to both bit errors and phase errors. Define a pair $e_0 = (e_0, \dots, e_{n-1})$, $e_1 = (e'_0, \dots, e'_{n-1}) \in \mathbb{F}_2^n$ of n -dimensional vectors such that for $0 \leq i \leq n-k-1$, $e_i = 1$ if a phase error occurred on q_i and $e_i = 0$ otherwise, such that for $n-k \leq i \leq n-1$, $e_i = 1$ if a bit error occurred on q_{i+n-k} and $e_i = 0$ otherwise, and such that for $0 \leq i \leq n-1$, $e'_i = 1$ if a phase error occurred on q_{i+n-k} and $e'_i = 0$ otherwise. Let H be a parity-check matrix of $\mathcal{C}$ in standard form. There exists a $[[2n-k, k]]$ quantum error-correcting code that allows for retrieving classical information about quantum errors in the form of a pair $s_0, s_1 \in \mathbb{F}_2^{n-k}$ of $(n-k)$ -dimensional vectors such that $s_0 = He_0^T$ and $s_1 = He_1^T$.

Note that the binary vectors e_0 and e_1 in the theorem above completely specify what type of quantum error occurred on which qubit. The correspondence between each bit of the error vectors e_0 , e_1 and the type and location of each quantum error is shown in Figure 1.

The most important fact regarding Theorem 2.2 is that because H is a parity-check matrix of a linear code of minimum distance d , we can correctly infer e_0 and e_1 from the syndromes s_0 and s_1 , which are He_0^T and He_1^T respectively, if the weights of e_0 and e_1 are both less than or equal to $\lfloor \frac{d-1}{2} \rfloor$. This implies that the positions of all bit errors and phase errors can be identified if the number of physical qubits that suffer bit errors, phase errors or both is at most $\lfloor \frac{d-1}{2} \rfloor$.

One restriction on employing Theorem 2.2 is that the parity-check matrix H of a linear $[n, k, d]$ code for quantum error correction must be in standard form

$$H = \begin{bmatrix} I & A \end{bmatrix}$$

for some $(n-k) \times k$ matrix A over $\mathbb{F}_2$, where I is the $(n-k) \times (n-k)$ identity matrix. Another restriction is that the linear code must be of high rate to keep the number $2(n-k)$ of less noisy auxiliary qubits small.

3 Application of Combinatorial Designs

As we have seen in the previous section, the framework for designing quantum error-correcting codes allows us to directly take advantage of any linear codes and any of their decoding methods that infer errors from syndromes of the form He^T , where e is a binary vector specifying the error position and H is a parity-check matrix. Among the known decoding methods based on syndromes, the one that infers e through iterative decoding such as the sum-product algorithm has been shown to achieve excellent error correction performance and very low decoding complexity in classical coding theory [36]. In particular,

linear codes whose parity-check matrices have relatively few nonzero entries tend to perform well when decoded by the sum-product algorithm. Such linear codes with sparse parity-check matrices suited for iterative decoding are called *low-density parity-check* (LDPC) codes. In this section we give a brief summary of how combinatorial designs have been applied to LDPC codes in the context of quantum error correction assisted by less noisy qubits.

The *Tanner graph* of an $m \times n$ parity-check matrix H is the bipartite graph consisting of n *bit vertices* indexed by bits of the corresponding code and m *parity-check vertices* indexed by parity-check equations defined by H , where an edge joins a bit vertex to a parity-check vertex if the bit is checked by the corresponding parity-check equation. An *i-cycle* in a graph is a sequence of $i + 1$ connected vertices which starts and ends at the same vertex in the graph and contains no other vertices more than once. In the language of matrices, a 4-cycle in a Tanner graph is equivalent to a 2×2 all-one submatrix in a parity-check matrix while a 6-cycle forms a 3×3 submatrix in which each row and column has exactly two ones. The *girth* of a parity-check matrix is the length of a shortest cycle in the corresponding Tanner graph. Since a Tanner graph is bipartite, its girth must be an even number. When it is clear from context which parity-check matrix is considered, we may speak of the “girth of an LDPC code.” It has empirically been shown that very short cycles tend to increase the probability of decoding failure when the sum-product algorithm is employed. In particular, 4-cycles have a particularly noticeable negative effect on error correction performance [30]. For this reason, it is desirable for the girth of a parity-check matrix to be strictly larger than 4. We consider the kind of combinatorial design that constitutes a rich source of such matrices.

Let K be a set of positive integers. A *pairwise balanced design* of order v and index 1 with block sizes from K , denoted by $\text{PBD}(v, K, 1)$, is an ordered pair $(V, \mathcal{B})$, where V is a nonempty finite set of v elements, called *points*, and $\mathcal{B}$ is a set of subsets of V , called *blocks*, that satisfies the following two conditions:

- (i) each unordered pair of distinct elements of V appears in exactly one block of $\mathcal{B}$,
- (ii) for every $B \in \mathcal{B}$ the cardinality $|B| \in K$.

When K is a singleton $\{\mu\}$, the PBD is called a *Steiner 2-design* of order v and block size μ , and is denoted by $S(2, \mu, v)$. A simple counting argument proves that the number of blocks in an $S(2, \mu, v)$ is exactly $\frac{v(v-1)}{\mu(\mu-1)}$. A PBD of order v is *trivial* if it has no blocks or consists of only one block of size v .

Let $\alpha(K) = \gcd\{\mu - 1 \mid \mu \in K\}$ and $\beta(K) = \gcd\{\mu(\mu - 1) \mid \mu \in K\}$. Necessary conditions for the existence of a $\text{PBD}(v, K, 1)$ are $v - 1 \equiv 0 \pmod{\alpha(K)}$ and $v(v - 1) \equiv 0 \pmod{\beta(K)}$ [4]. These conditions were proved to be asymptotically sufficient.

Theorem 3.1 (Wilson [50]) *There exists a constant v_K such that for every $v > v_K$ satisfying $v - 1 \equiv 0 \pmod{\alpha(K)}$ and $v(v - 1) \equiv 0 \pmod{\beta(K)}$ there exists a $\text{PBD}(v, K, 1)$.*

An *incidence matrix* of a $\text{PBD}(V, \mathcal{B})$ with $|V| = v$ and $|\mathcal{B}| = b$ is a binary $v \times b$ matrix $H = (h_{i,j})$ with rows indexed by points, columns indexed by blocks, and $h_{i,j} = 1$ if the i th point is contained in the j th block, and $h_{i,j} = 0$ otherwise. It is known that incidence matrices of PBDs of index 1 are generally good candidates of parity-check matrices of LDPC codes of high rate because of their good error tolerance at relatively short lengths [46, 1, 31, 33]. In fact, a simple observation shows that incidence matrices of PBDs of index 1 generally do not contain 4-cycles and form parity-check matrices of girth 6 [17]. Furthermore, such parity-check matrices are extreme in terms of information rates in the following sense.

Proposition 3.2 [17] *Let H be a parity-check matrix in standard form that forms an incidence matrix of a nontrivial PBD of order $n - k$ and index 1. Any parity-check matrix of the same size, same column weight distribution, same or higher girth, and same or higher minimum distance as H is of the same or lower rate as H .*

Proposition 3.3 [17]. *Let H be a parity-check matrix in standard form that forms an incidence matrix of a nontrivial PBD of order $n - k$ and index 1. Take a nonzero vector $c \in \mathbb{F}_2^{n-k}$ of dimension $n - k$. Adding c^T to H as a column results in a parity-check matrix of an LDPC code of minimum distance 2 or girth 4.*

Roughly speaking, the above propositions state that H in standard form forming an incidence matrix of a PBD of index 1 achieves the largest possible number of information bits among all parity-check matrices in standard form with the same number of check equations and the same degree distribution under the

constraint that the associated LDPC code must be of minimum distance larger than 2 and girth greater than 4.

Recall that Theorem 2.2 has the explicit restriction on the structure of a parity-check matrix H of a linear $[n, k, d]$ code, that is, it must be in standard form

$$H = \begin{bmatrix} I & A \end{bmatrix}$$

for some $(n - k) \times k$ matrix A over $\mathbb{F}_2$, where I is the $(n - k) \times (n - k)$ identity matrix. Therefore, if one wishes to achieve the highest possible information rates for an LDPC code of girth 6 or larger, H as a whole must form an incidence matrix of a PBD of index 1. The following proposition allows us to only consider the part A in this regard.

Proposition 3.4 [17] *Let $H = \begin{bmatrix} I & A \end{bmatrix}$ be a parity-check matrix of a linear code of length n and dimension k in standard form. H is an incidence matrix of a PBD of index 1 if and only if the $(n - k) \times k$ matrix A is an incidence matrix of a PBD of index 1 containing no block of size 1.*

In view of the above proposition, we would like to find incidence matrices A of PBDs without blocks of size 1 that do not contain or produce undesirable structures that hamper decoding when combined with the identity matrices to obtain valid parity-check matrices $H = \begin{bmatrix} I & A \end{bmatrix}$ for Theorem 2.2. Because an LDPC code is a linear code equipped with a particular decoding algorithm, all else being equal, it is desirable to have a larger minimum distance. While the sum-product algorithm is generally less sensitive to the minimum distance than other simple decoding methods, this parameter is especially important to a code of very high rate because its very large dimension dictates that the minimum distance be small compared to the length. The following proposition concerns with the number of short cycles and minimum distances of parity-check matrices based on incidence matrices of PBDs together with columns of weight 1.

Proposition 3.5 [17] *Let A be an $(n - k) \times k$ incidence matrix of a nontrivial $\text{PBD}(n - k, K, 1)$ with $1 \notin K$. Then the binary matrix $H = \begin{bmatrix} I & A \end{bmatrix}$ is a parity-check matrix of a linear $[n, k, d]$ code in standard form whose girth is 6 and minimum distance $d = 1 + \min\{\mu \mid \mu \in K\}$.*

Because the minimum distance of our code is $1 + \min\{\mu \mid \mu \in K\}$, increasing the smallest block size improves the minimum distance. However, because a block of size x contains $\binom{x}{2}$ pairs, a block of larger size contains more pairs of points. Since avoiding 4-cycles while achieving the highest possible information rate is equivalent to packing as many different pairs of points as possible in a set of blocks while including no pair of points more than once, generally speaking, increasing block sizes lowers the achievable information rate. For this reason, we consider the case when the column weights of the matrix A are uniform. This means that K is a singleton $\{\mu\}$, that is, the corresponding $\text{PBD}(n - k, K, 1)$ forms a Steiner 2-design $S(2, \mu, n - k)$. As stated earlier, an $S(2, \mu, v)$ contains exactly $\frac{v(v-1)}{\mu(\mu-1)}$ blocks. Thus, the corresponding code can achieve quite a large dimension at a moderate length.

Proposition 3.6 *Let A be an incidence matrix of an $S(2, \mu, v)$ and I a $v \times v$ identity matrix. A parity-check matrix $H = \begin{bmatrix} I & A \end{bmatrix}$ defines an LDPC code of length $\frac{v(v-1)}{\mu(\mu-1)} + v$, dimension $\frac{v(v-1)}{\mu(\mu-1)}$, girth 6, and minimum distance $\mu + 1$.*

As can be seen from the above proposition, the information rates of LDPC codes defined by incidence matrices of $S(2, \mu, v)$ s become close to 1 very quickly as v goes to infinity. Theorem 2.2 shows that the corresponding quantum error-correcting codes assisted by less noisy qubits inherit this characteristic.

Theorem 3.7 *Let A be an incidence matrix of an $S(2, \mu, v)$ and I a $v \times v$ identity matrix. There exists a $\left\lceil \left[\frac{v(v-1)}{\mu(\mu-1)} + 2v, \frac{v(v-1)}{\mu(\mu-1)} \right] \right\rceil$ quantum error-correcting code that identifies the types and locations of quantum errors through the LDPC code defined by the parity-check matrix $H = \begin{bmatrix} I & A \end{bmatrix}$ under the assumption that a fixed set of $2(n - k)$ physical qubits may experience phase errors but no bit errors.*

A good strategy to improve the error correction performance under the sum-product algorithm is to decrease the number of structures that are responsible for dominating errors. While joining the identity matrix and the incidence matrix A of a Steiner 2-design of block size μ necessarily results in an LDPC code

of minimum distance $\mu + 1$, it is desirable for the linear code defined by A to have a larger minimum distance because it eliminates dominating sources of errors to an extent.

It is trivial that the minimum distance of a linear code whose parity-check matrix is an incidence matrix of an $S(2, \mu, v)$ is lower bounded by $\mu + 1$. To investigate the minimum distances of linear codes based on Steiner 2-designs further, we define some combinatorial notions.

A configuration $\mathcal{C}$ in an $S(2, \mu, k)$, $(V, \mathcal{B})$, is a subset $\mathcal{C} \subseteq \mathcal{B}$ of the block set. The set of points appearing in at least one block of a configuration $\mathcal{C}$ is denoted by $V(\mathcal{C})$. Two configurations $\mathcal{C}$ and $\mathcal{C}'$ are *isomorphic* if there exists a bijection $\phi : V(\mathcal{C}) \rightarrow V(\mathcal{C}')$ such that for each block $B \in \mathcal{C}$, the image $\phi(B)$ is a block in $\mathcal{C}'$. When $|\mathcal{C}| = i$, a configuration $\mathcal{C}$ is called an *i-configuration*. A configuration $\mathcal{C}$ is *even* if for every point a appearing in $\mathcal{C}$ the number $|\{B \mid a \in B \in \mathcal{C}\}|$ of blocks containing a is even. An $S(2, \mu, v)$ is *r-even-free* if for every integer i satisfying $1 \leq i \leq r$ it contains no even i -configurations. Because the minimum distance of a linear code is the size of a smallest linearly dependent set of columns in its parity-check matrix, the minimum distance of a linear code based on an incidence matrix A of a Steiner 2-design is determined by its even-freeness.

Proposition 3.8 *The minimum distance of a linear code whose parity-check matrix forms an incidence matrix of a Steiner 2-design is d if and only if the corresponding Steiner 2-design is $(d - 1)$ -even-free but not d -even-free.*

By definition every r -even-free $S(2, \mu, v)$, $r \geq 2$, is also $(r - 1)$ -even-free. If μ is odd, a simple double counting argument proves that an r -even-free $S(2, \mu, v)$ with r even is also $(r + 1)$ -even-free. Because a Steiner 2-design is a linear space, every $S(2, \mu, v)$ is μ -even-free.

A nontrivial $S(2, \mu, v)$ may or may not be $(\mu + 1)$ -even-free. For each $\mu \geq 2$, an even $(\mu + 1)$ -configuration that may be in $S(2, \mu, v)$ s is unique up to isomorphism; they are the dual of the complete graph on $\mu + 1$ vertices. For instance, for the case when $\mu = 3$, up to isomorphism, there exists only one possible even 4-configuration, called the *Pasch* configuration, which can be written by six points and four blocks:

$$\{\{a, b, c\}, \{a, d, e\}, \{f, b, d\}, \{f, c, e\}\}.$$

The unique possible even $(\mu + 1)$ -configurations for $\mu \geq 4$ are sometimes called the *generalized Pasch* configurations in the coding theory literature. Since they are the smallest and unique, an $S(2, \mu, v)$ is $(\mu + 1)$ -even-free if and only if it contains no Pasch configurations for $\mu = 3$ and no generalized Pasch configurations for $\mu \geq 4$.

The following is the tightest known bounds on the maximum even-freeness of an $S(2, 3, v)$ stated in terms of minimum distance.

Theorem 3.9 ([16]) *The minimum distance d of a linear code whose parity-check matrix forms an incidence matrix of a nontrivial $S(2, 3, v)$ satisfies the inequalities $4 \leq d \leq 8$.*

The problem of avoiding Pasch configurations has long been investigated in various contexts in design theory. The spectrum of those orders v for which a 5-even-free $S(2, 3, v)$ exists was completely determined in 2000 [22].

Theorem 3.10 ([22]) *There exists a 5-even-free $S(2, 3, v)$ if and only if $v \equiv 1, 3 \pmod{6}$ except $v = 7, 13$.*

While it is enough to attain $(\mu + 1)$ -even-freeness in the right portion A of our parity-check matrix $H = \begin{bmatrix} I & A \end{bmatrix}$ for achieve the goal of reducing the number of codewords of the smallest weight, if one would like even higher even-freeness, it is required to construct an $S(2, 3, v)$ that simultaneously avoids Pasch and two more even configurations, namely the *grids*

$$\{\{a, b, c\}, \{d, e, f\}, \{g, h, i\}, \{a, d, g\}, \{b, e, h\}, \{c, f, i\}\}$$

and *double triangles*

$$\{\{a, b, c\}, \{a, d, e\}, \{b, f, g\}, \{c, h, e\}, \{d, g, i\}, \{f, h, i\}\}.$$

Unfortunately, while there exist infinitely many $S(2, 3, v)$ s avoiding both Pasch and double triangle configurations [9], at the time of writing, no nontrivial examples avoiding grids, let alone 6-even-free $S(2, 3, v)$ s, are known [19]. If a nontrivial $S(2, 3, v)$ that simultaneously avoids the three even configurations exists, it is automatically 7-even-free and hence attains the upper bound given in Theorem 3.9.

It is tempting to look for similar theorems on the even-freeness of $S(2, \mu, \nu)$ s for all $\mu \geq 4$. Unfortunately, it appears very difficult to obtain equally tight bounds and/or complete existence results for relatively high even-freeness for general block size μ . In fact, no nontrivial upper bounds are known on the even-freeness of $S(2, \mu, \nu)$ s with large μ in general. To the best of the author's knowledge, the following is the only fairly general bound on the even-freeness of $S(2, \mu, \nu)$ s.

Theorem 3.11 ([13]) *If an abelian group acts transitively on the points of a nontrivial τ -even-free $S(2, \mu, \nu)$ with $\nu > \mu(\mu - 1) + 1$, then $\tau \leq 2\mu - 1$.*

The above bound is best possible in the sense that there are infinitely many $(2\mu - 1)$ -even-free $S(2, \mu, \nu)$ s. To give such infinite series, we use the points and lines of finite geometries.

The *affine geometry* $AG(m, q)$ of dimension m over $\mathbb{F}_q$ is defined as a finite geometry in which the *points* are the vectors in $\mathbb{F}_q^m$ and the *i -dimensional affine subspaces* are the i -dimensional vector subspaces of $\mathbb{F}_q^m$ and their cosets. The points and 1-dimensional affine subspaces of $AG(m, q)$ form the points and blocks of an $S(2, q, q^m)$ [4].

Theorem 3.12 ([38]) *For any odd prime power q and positive integer $m \geq 2$ the points and 1-dimensional affine subspaces of $AG(m, q)$ form a $(2q - 1)$ -even-free $S(2, q, q^m)$ which is not $2q$ -even-free.*

Affine geometries are not the only known nontrivial $S(2, \mu, \nu)$ s that attain the upper bound given in Theorem 3.11. The *projective geometry* $PG(m, q)$ of dimension m over $\mathbb{F}_q$ is a finite geometry whose *points* and *i -dimensional subspaces* are the 1-dimensional vector subspaces and the $(i + 1)$ -dimensional vector subspaces of $\mathbb{F}_q^{m+1}$ respectively. The points and 1-dimensional subspaces of $PG(m, q)$ form the points and blocks of an $S(2, q + 1, \frac{q^{m+1}-1}{q-1})$.

Theorem 3.13 ([15]) *For any odd prime power q and positive integer $m \geq 3$ the points and 1-dimensional subspaces of $PG(m, q)$ form a $(2q + 1)$ -even-free $S(2, q + 1, \frac{q^{m+1}-1}{q-1})$ which is not $(2q + 2)$ -even-free.*

Recently, the author gave a combinatorial construction for $(\mu + 1)$ -even-free $S(2, \mu, \nu)$ s [13]. The construction technique recursively combines a $(\mu + 1)$ -even-free $S(2, \mu, \nu)$ and another $(\mu + 1)$ -even-free $S(2, \mu, \nu)$ with a particular algebraic property by using a specially designed combinatorial matrix to generate a larger $(\mu + 1)$ -even-free $S(2, \mu, \nu w)$. As far as the author is aware, no constructions for $S(2, \mu, \nu)$ s with even-freeness higher than or equal to $\mu + 1$ are known except the finite geometric and recursive ones.

From the viewpoint of quantum error correction assisted by less noisy qubits, the highly even-free $S(2, \mu, \nu)$ s based on affine and projective geometries have an additional appealing property. As described in the previous section, our auxiliary qubits are assumed to be engineered more reliably than the rest so that only phase errors can occur. Hence, it would be natural to assume that those phase errors that may still occur on these qubits manifest less frequently than bit errors and phase errors on the other qubits. This slightly more optimistic assumption translates into the premise that the probability that an error occurs on a fixed check bit, which corresponds to a column of the $(n - k) \times (n - k)$ identity matrix I in $H = \begin{bmatrix} I & A \end{bmatrix}$, is smaller than that on a fixed information bit corresponding to a column of A (see [17]). The following theorem shows that highly even-free $S(2, \mu, \nu)$ s from finite geometries can effectively take advantage of this additional assumption.

Theorem 3.14 ([38, 45]) *Let q be an odd prime power and $m \geq 2$ an integer greater than or equal to 2. Define $(V, \mathcal{B})$ to be the $(2q - 1)$ -even-free $S(2, q, q^m)$ formed by the points and 1-dimensional affine subspaces of $AG(m, q)$. For any nonempty configuration $\mathcal{C} \subset \mathcal{B}$ of size $|\mathcal{C}| \leq 2q - 1$, it holds that*

$$|\mathcal{C}| + \text{odd}(\mathcal{C}) \geq 2q,$$

where $\text{odd}(\mathcal{C})$ is the number of points $v \in V$ contained in an exactly odd number of blocks in $\mathcal{C}$.

Theorem 3.15 ([45]) *Let q be an odd prime power and $m \geq 2$ an integer greater than or equal to 3. Define $(V, \mathcal{B})$ to be the $(2q + 1)$ -even-free $S(2, q + 1, \frac{q^{m+1}-1}{q-1})$ formed by the points and 1-dimensional subspaces of $PG(m, q)$. For any nonempty configuration $\mathcal{C} \subset \mathcal{B}$ of size $|\mathcal{C}| \leq 2q + 1$, it holds that*

$$|\mathcal{C}| + \text{odd}(\mathcal{C}) \geq 2q + 2.$$

The point of the above two theorems is that a linear code of minimum distance $\mu + 1$ defined by a parity-check matrix $H = \begin{bmatrix} I & A \end{bmatrix}$ with A being an incidence matrix of a finite geometric $S(2, \mu, v)$ would perform better if check bits suffer from errors much less likely than information bits. This is because the code is almost of minimum distance 2μ in the sense that there is only one type of the smallest weight codeword, which is the one that consists of one information bit and the corresponding μ check bits. An error of this type would be unlikely to occur under the additional assumption that the more reliable qubits have a sufficiently smaller error probability.

4 Concluding Remarks

We have seen one example case where design theory and coding theory play the central role in quantum error correction. Because there are an abundance of results on PBDs in combinatorics, the framework we reviewed can import various theoretical results directly from design theory and coding theory.

One notable result we did not mention is that, as expected from the theoretical analysis presented in this note, quantum LDPC codes from combinatorial designs were shown to have good error performance through simulations [17]. It is also worth noting that the effect of almost doubled minimum distances of finite geometric codes would very likely be favorable across many different decoding methods and algorithms. For instance, in the case of iterative decoding, a nonempty set of bit vertices in a Tanner graph that are not correct after l iterations for all $l \geq l_c$ for some absolute constant l_c is called a *trapping set* [39]. To improve the error floor and slope of the block error rate curve, it is desirable for a parity-check matrix to avoid small trapping sets [29]. While the set of small trapping sets generally varies from algorithm to algorithm and notoriously difficult to identify, codewords of very small weight are surely among them. We hope that design theory and coding theory will find more and more applications and keep playing important roles in quantum information science.

References

- [1] B. Ammar, B. Honary, Y. Kou, J. Xu, and S. Lin, *Construction of low-density parity-check codes based on balanced incomplete block designs*, IEEE Trans. Inf. Theory **50** (2004), 1257–1268.
- [2] T. Aoki et al., *Quantum error correction beyond qubits*, Nature Physics **5** (2009), 541–546.
- [3] S. Bartz et al., *Demonstrating elements of measurement-based quantum error correction*, e-print arXiv:1308.5209, 2013.
- [4] T. Beth, D. Jungnickel, and H. Lenz, *Design Theory*, Cambridge Univ. Press, Cambridge, 1999.
- [5] N. Boulant, L. Viola, E. Fortunato, and D. G. Cory, *Experimental implementation of a concatenated quantum error-correcting code*, Phys. Rev. Lett. **94** (2005), 130501.
- [6] T. A. Brun, I. Devetak, and M.-H. Hsieh, *Correcting quantum errors with entanglement*, Science **314** (2006), 436–439.
- [7] A. R. Calderbank, E. M. Rains, P. W. Shor, and N. J. A. Sloane, *Quantum error correction via codes over $GF(4)$* , IEEE Trans. Inf. Theory **44** (1998), 1369–1387.
- [8] J. Chiaverini et al., *Realization of quantum error correction*, Nature **432** (2004), 602–605.
- [9] C. J. Colbourn and Y. Fujiwara, *Small stopping sets in Steiner triple systems*, Cryptogr. Commun. **1** (2009), 31–46.
- [10] D. G. Cory et al., *Experimental quantum error correction*, Phys. Rev. Lett. **81** (1998), 2152–2155.
- [11] I. B. Djordjevic, *Photonic entanglement-assisted quantum low-density parity-check encoders and decoders*, Optics Lett. **35** (2010), 1464–1466.

- [12] Y. Dong, X. Deng, M. Jiang, Q. Chen, and S. Yu, *Entanglement-enhanced quantum error-correcting codes*, Phys. Rev. A **79** (2009), 042342.
- [13] Y. Fujiwara, *Even-freeness of cyclic 2-designs*, e-print arXiv:1210.7516, 2012.
- [14] ———, *Quantum error correction via less noisy qubits*, Phys. Rev. Lett. **110** (2013), 170501.
- [15] Y. Fujiwara, D. Clark, P. Vandendriessche, M. De Boeck, and V. D. Tonchev, *Entanglement-assisted quantum low-density parity-check codes*, Phys. Rev. A **82** (2010), 042338.
- [16] Y. Fujiwara and C. J. Colbourn, *A combinatorial approach to X-tolerant compaction circuits*, IEEE Trans. Inf. Theory **56** (2010), 3196–3206.
- [17] Y. Fujiwara, A. Gruner, and P. Vandendriessche, *High-rate quantum low-density parity-check codes assisted by reliable qubits*, e-print arXiv:1309.5587, 2013.
- [18] Y. Fujiwara and V. D. Tonchev, *A characterization of entanglement-assisted quantum low-density parity-check codes*, IEEE Trans. Inf. Theory **59** (2013), 3347–3353.
- [19] Z. Füredi and M. Ruszinkó, *Uniform hypergraphs containing no grids*, Adv. Math. **240** (2013), 302–324.
- [20] D. Gottesman, *Stabilizer codes and quantum error correction*, Ph.D. thesis, California Institute of Technology, 1997.
- [21] ———, *An introduction to quantum error correction and fault-tolerant quantum computation*, Quantum Information Science and Its Contributions to Mathematics (Providence, Rhode Island) (S. J. Lomonaco, Jr., ed.), Proceedings of Symposia in Applied Mathematics, vol. 68, American Mathematical Society, 2010, pp. 13–58.
- [22] M. J. Grannell, T. S. Griggs, and C. A. Whitehead, *The resolution of the anti-Pasch conjecture*, J. Combin. Des. **8** (2000), 300–309.
- [23] L. Guo and R. Li, *Linear plotkin bound for entanglement-assisted quantum codes*, Phys. Rev. A **87** (2013), 032309.
- [24] M.-H. Hsieh, T. A. Brun, and I. Devetak, *Entanglement-assisted quantum quasicyclic low-density parity-check codes*, Phys. Rev. A **79** (2009), 032340.
- [25] M.-H. Hsieh, I. Devetak, and T. A. Brun, *General entanglement-assisted quantum error-correcting codes*, Phys. Rev. A **76** (2007), 062313.
- [26] M.-H. Hsieh, W.-T. Yen, and L.-Y. Hsu, *High performance entanglement-assisted quantum LDPC codes need little entanglement*, IEEE Trans. Inf. Theory **57** (2011), 1761–1769.
- [27] W. C. Huffman and V. Pless, *Fundamentals of error-correcting codes*, Cambridge Univ. Press, Cambridge, 2003.
- [28] L. Ioffe and M. Mézard, *Asymmetric quantum error-correcting codes*, Phys. Rev. A **75** (2007), 032345.
- [29] M. Ivković, S. K. Chilappagari, and B. Vasić, *Eliminating trapping sets in low-density parity-check codes by using tanner graph covers*, IEEE Trans. Inf. Theory **54** (2008), 3763–3768.
- [30] S. J. Johnson, *Iterative error correction: Turbo, low-density parity-check and repeat-accumulate codes*, Cambridge Univ. Press, New York, 2010.
- [31] S. J. Johnson and S. R. Weller, *Resolvable 2-designs for regular low-density parity-check codes*, IEEE Trans. Commun. **51** (2003), 1413–1419.
- [32] E. Knill, R. Laflamme, R. Martinez, and C. Negrevergne, *Benchmarking quantum computers: the five-qubit error correcting code*, Phys. Rev. Lett. **86** (2001), 5811–5814.

- [33] Y. Kou, S. Lin, and M. P. C. Fossorier, *Low-density parity-check codes based on finite geometries: A rediscovery and new results*, IEEE Trans. Inf. Theory **47** (2001), 2711–2736.
- [34] C.-Y. Lai and T. A. Brun, *Entanglement-assisted quantum error-correcting codes with imperfect ebits*, Phys. Rev. A **86** (2012), 032319.
- [35] C.-Y. Lu et al., *Experimental quantum coding against qubit loss error*, Proc. Natl. Acad. Sci. **105** (2008), 11050–11054.
- [36] D. J. C. MacKay, *Information Theory, Inference, and Learning Algorithms*, Cambridge University Press, Cambridge, 2003.
- [37] O. Moussa, J. Baugh, C. A. Ryan, and R. Laflamme, *Demonstration of sufficient control for two rounds of quantum error correction in a solid state ensemble quantum information processor*, Phys. Rev. Lett. **107** (2011), 160501.
- [38] M. Müller and M. Jimbo, *Erasur-resilient codes from affine spaces*, Discrete Appl. Math. **143** (2004), 292–297.
- [39] D. V. Nguyen, S. K. Chilappagari, M. W. Marcellin, and B. Vasić, *On the construction of structured LDPC codes free of small trapping sets*, IEEE Trans. Inf. Theory **58** (2012), 2280–2302.
- [40] M. A. Nielsen and I. L. Chuang, *Quantum Computation and Quantum Information*, Cambridge Univ. Press, New York, 2000.
- [41] M. D. Reed et al., *Realization of three-qubit quantum error correction with superconducting circuits*, Nature **482** (2012), 382–385.
- [42] P. Schindler et al., *Experimental repetitive quantum error correction*, Science **332** (2011), 1059–1061.
- [43] P. W. Shor, *Scheme for reducing decoherence in quantum computer memory*, Phys. Rev. A **52** (1995), R2493–R2496.
- [44] A. M. Steane, *Error-correcting codes in quantum theory*, Phys. Rev. Lett. **77** (1996), 793–797.
- [45] P. Vandendriessche, *On small line sets with few odd points*, preprint, 2013.
- [46] B. Vasić and O. Milenkovic, *Combinatorial constructions of low-density parity-check codes for iterative decoding*, IEEE Trans. Inf. Theory **50** (2004), 1156–1176.
- [47] M. M. Wilde and T. A. Brun, *Optimal entanglement formulas for entanglement-assisted quantum coding*, Phys. Rev. A **77** (2008), 064302.
- [48] ———, *Entanglement-assisted quantum convolutional coding*, Phys. Rev. A **81** (2010), 042333.
- [49] ———, *Quantum convolutional coding with shared entanglement: general structure*, Quantum Inf. Processing **9** (2010), 509–540.
- [50] R. M. Wilson, *An existence theory for pairwise balanced designs, II: the structure of PBD-closed set and the existence conjecture*, J. Combin. Theory Ser. A **13** (1972), 246–273.
- [51] X.-C. Yao et al., *Experimental demonstration of topological error correction*, Nature **482** (2012), 489–494.
- [52] J. Zhang, M. Grassl, B. Zeng, and R. Laflamme, *Experimental implementation of a codeword-stabilized quantum code*, Phys. Rev. A **85** (2012), 062312.
- [53] J. Zhang, R. Laflamme, and D. Suter, *Experimental implementation of encoded logical qubit operations in a perfect quantum error correcting code*, Phys. Rev. Lett. **109** (2012), 100503.