



Декодирование одиночных пакетов ошибок по минимуму длины пакета на основе информационных совокупностей

М. Н. Исаева^а, старший преподаватель, orcid.org/0009-0007-6228-0617, imn@guap.ru

^аСанкт-Петербургский государственный университет аэрокосмического приборостроения, Б. Морская ул., 67, Санкт-Петербург, 190000, РФ

Введение: для улучшения характеристик систем связи и хранения информации необходимо учитывать специфику шума, которая во многих практических ситуациях приводит к эффекту пакетирования ошибок. Декодирование независимых ошибок по информационным совокупностям является лучшим методом декодирования случайных линейных кодов, но имеет экспоненциальную сложность. **Цель:** анализ корректирующей способности и вычислительной сложности алгоритмов на основе информационных совокупностей для исправления одиночных пакетов ошибок. **Результаты:** введен критерий декодирования одиночных пакетов с использованием информационных совокупностей на основе минимума длины пакета. Доказано, что такой критерий обеспечивает исправление всех пакетов ошибок в пределах корректирующей способности линейного кода. Рассмотрены вопросы уменьшения количества необходимых для правильного декодирования информационных совокупностей, в целях чего исследованы методики построения множества информационных совокупностей ограниченного диаметра, называемых плотными. Проведен анализ, показывающий возможность снизить мощность множества информационных совокупностей с $O(n)$ до $O(1)$. Оценена вероятность построения множества и сформулирована оптимизационная задача по максимизации этой вероятности. **Практическая значимость:** разработанный декодер одиночных пакетов ошибок на основе критерия минимальной длины пакета, а также методика построения плотных информационных совокупностей мощностью $O(1)$ гарантируют реализацию корректирующей способности случайных линейных кодов при исправлении одиночных пакетов ошибок. **Обсуждение:** полученные результаты подтверждают, что задача исправления однократных пакетов ошибок случайными линейными кодами является полиномиальной. Вместе с тем рассмотренные методы позволяют получать множества информационных совокупностей с достаточно высокой вероятностью лишь для кодов с не очень большой скоростью, а при использовании других, более практических классов кодов, вероятности нахождения множества информационных совокупностей могут достаточно сильно измениться, что является направлением дальнейших исследований.

Ключевые слова — декодирование по информационным совокупностям, граница Рейгера, исправление пакетов ошибок, плотные информационные совокупности, критерии декодирования.

Для цитирования: Исаева М. Н. Декодирование одиночных пакетов ошибок по минимуму длины пакета на основе информационных совокупностей. *Информационно-управляющие системы*, 2025, № 2, с. 68–77. doi:10.31799/1684-8853-2025-2-68-77, EDN: MANCAY

For citation: Isaeva M. N. Decoding of single error bursts using minimal burst length criteria and information sets. *Informatsionno-upravliaiushchie sistemy* [Information and Control Systems], 2025, no. 2, pp. 68–77 (In Russian). doi:10.31799/1684-8853-2025-2-68-77, EDN: MANCAY

Введение

В современном мире информационные потоки внедрились во все сферы деятельности человека. Большую значимость имеет передача информации (например, в двоичном виде) по каналам связи, в которых могут возникать различные помехи, в том числе из-за особенностей среды передачи. Актуальной задачей является исправление ошибок, возникающих в передаваемых битовых последовательностях. Для сокращения числа ошибок используются помехоустойчивые коды, для которых необходимо разрабатывать специальные методы декодирования, позволявшие бы исправлять значительное число ошибок (обеспечивать малую вероятность ошибки декодирования) при условии ограниченных вычислительных и емкостных ресурсов.

В большинстве кодов при декодировании ошибки рассматриваются как независимые события, однако в реальности ошибки на выходе канала часто группируются в пакеты ошибок, например из-за особенностей среды передачи. Исправление таких пакетов является менее исследованной задачей в теории кодирования, чем исправление независимых ошибок, тем не менее известно, что учет особенностей векторов ошибок позволяет строить схемы кодирования и декодирования, теоретически позволяющие достичь лучших характеристик системы связи. В литературе, как правило, при решении задачи борьбы с ошибками в каналах с памятью применяется либо построение специальных кодов и традиционных декодеров [1–7], либо, достаточно редко — специальных декодеров для специфических моделей каналов связи и классических схем

кодирования [8–13]. В то же время кажется, что такая задача должна решаться совместно алгоритмами обработки информации на передатчике и приемнике.

В данной статье исследуется декодирование для исправления одиночных пакетов ошибок для самого общего класса кодов — линейных кодов, являющихся линейными векторными подпространствами и задающимися либо своим базисом (порождающей матрицей), либо базисом ортогонального пространства (проверочной матрицей). Случайные линейные коды обладают хорошей корректирующей способностью как в каналах с независимыми ошибками, так и в каналах с памятью, хотя их использование ограничено отсутствием структуры кода, которую можно использовать для повышения вычислительной эффективности процедур кодирования и декодирования.

Исправление одиночных пакетов ошибок линейными кодами изучалось, например, в статье [12] с помощью декодера со скользящим окном, однако сложность данного декодера для линейных кодов в общем случае оценивается авторами как экспоненциальная. В [14] рассматривается исправление пакетов ошибок на основе информационных совокупностей и исследуется способность некоторых кодов (с малыми параметрами) исправлять ошибки сверх корректирующей способности. В [15] анализируется экспонента сложности декодирования «почти» по максимуму правдоподобия при применении информационных совокупностей в каналах с двумя состояниями при наличии у декодера информации о состояниях канала или надежностях символов.

В данной статье предложена вероятностная методика построения множеств с малым числом совокупностей при максимизации вероятности построения множества для случайных линейных кодов.

Декодирование по информационным совокупностям для исправления пакетов ошибок

Линейный двоичный (n, k) -код — это k -мерное подпространство n -мерного линейного векторного пространства двоичных векторов. Любой базис G этого пространства является $(k \times n)$ -матрицей ранга k и называется порождающей матрицей кода. Матрица H размерности $(r \times n)$, где $r = n - k$, для которой $GH^T = 0$, является базисом ортогонального пространства размерности r и называется проверочной матрицей кода. Сообщение m , представленное вектором длиной k , кодируется в кодовое слово a с помощью преоб-

разования $mG = a$. Вектор $S = bH^T$ размерности r называется синдромом вектора b и является нулевым вектором тогда и только тогда, когда b — кодовое слово [16–18].

Способность кода исправлять независимые битовые ошибки (корректирующую способность кода) обычно связывают с параметром d_0 — минимальным расстоянием кода, равным минимуму из расстояний Хемминга между всеми парами кодовых слов. Код с расстоянием $d_0 = 2t + 1$ гарантирует исправление любой комбинации из не более чем t битовых ошибок, если декодирование заключается в поиске кодового слова, ближайшего к принятому из канала, где расстояние считается в метрике Хемминга. В действительности код может исправлять и значительное число комбинаций из более чем t ошибок, полный их список можно получить, если рассмотреть разбиение n -мерного векторного пространства по коду на смежные классы [14, 16]. Векторы одного смежного класса имеют одинаковый синдром (например, у всех кодовых слов он нулевой), и если выбрать из каждого смежного класса одного представителя, называемого лидером, и считать его возможным вектором ошибки, можно установить взаимно однозначное соответствие между синдромами и лидерами, что дает список ошибок, исправляемых кодом, и процедуру декодирования, использующую установленное соответствие, называемую синдромным декодированием и имеющую экспоненциальную сложность. Обычно в качестве лидеров смежных классов выбираются наиболее вероятные векторы в канале связи.

Стоит заметить, что как декодирование по минимуму расстояния (эквивалентное декодированию лидеров смежных классов), так и определение минимального расстояния кода являются NP-трудными задачами.

Если в векторе e длиной n первая ненулевая компонента расположена на позиции i_1 , а последняя ненулевая компонента — на позиции i_2 , то говорят, что вектор e образует однократный пакет длиной $b = i_2 - i_1 + 1$. Иногда может возникнуть неоднозначность в трактовке термина «пакет»: он может подразумевать как сам вектор e , так и подвектор на позициях от i_1 до i_2 (так называемых позициях пакета). Конкретный смысл этого термина обычно ясен из контекста.

Таким образом, все ненулевые компоненты вектора e находятся на позициях внутри интервала от i_1 до i_2 или на позициях пакета. Иногда пакетом длиной b считается также вектор e , все ненулевые компоненты которого расположены на позициях пакета, но сами позиции i_1 и (или) i_2 могут содержать нулевые элементы, мы будем в этом случае говорить, что e образует пакет длиной, не превышающей b .

Для определенности в дальнейшем будем считать, что позиции векторов нумеруются, начиная с единицы. Назовем позицию i_1 началом пакета, тогда, если мы рассматриваем пакеты, длина которых в точности равна b , начало пакета может принимать значения от 1 до $n - b + 1$.

В некоторых случаях, особенно при рассмотрении циклических кодов, вводится понятие циклического пакета, когда отсчет позиций пакета может происходить от i_2 до i_1 с переходом от позиции n снова к единице, однако в данной статье мы не рассматриваем циклические коды и циклические пакеты.

По аналогии с параметром t можно рассмотреть величину b_0 : если все пакеты, образующие пакеты длиной $b \leq b_0$, находятся в лидерах смежных классов, то любой пакет длиной, не превышающей b_0 , может быть исправлен. Исправление пакетов сверх корректирующей способности кода b_0 рассматривалось в работе [14] для кодов малых длин, где возможно построение списка лидеров смежных классов. Для линейного (n, k) -кода в соответствии с границей Рейгера справедливо $b_0 \leq (n - k)/2$.

Важно, что в отличие от параметра t , параметр b_0 для любого линейного кода может быть определен с полиномиальной сложностью [19], это дает основание надеяться, что и исправление одиночных пакетов может оказаться полиномиальной задачей, однако в данной статье мы ограничиваемся рассмотрением только исправления пакетов длиной до b_0 , предполагая, что для любого кода это значение известно.

Если линейный (n, k) -код не обладает никакой дополнительной структурой, то лучшим подходом к исправлению ошибок является декодирование по информационным совокупностям [14, 20]. Если ошибки независимы, количество требуемых информационных совокупностей экспоненциально, однако при исправлении пакетов их количество может быть значительно снижено [14].

Пусть $\mathbf{a} = (a_1, a_2, \dots, a_n)$ — кодовое слово. Если $\chi = \{1 \leq j_1 < j_2 < \dots < j_k \leq n\}$ — множество чисел, состоящее из сочетания (без повторений) k чисел от 1 до n , обозначим через $\mathbf{a}(\chi)$ подвектор $\mathbf{a}$, составленный из элементов на позициях с номерами из χ . Аналогично, если $\mathbf{A}$ — $(k \times n)$ -матрица, обозначим через $\mathbf{A}(\chi)$ $(k \times k)$ -подматрицу $\mathbf{A}$ из столбцов с номерами из χ . Множество χ называется информационной совокупностью, если любое кодовое слово $\mathbf{a}$ может быть однозначно восстановлено, $\mathbf{a}(\chi)$.

Является ли некоторое множество чисел χ информационной совокупностью, зависит не от конкретного кодового слова $\mathbf{a}$, а от кода. Определить это можно с помощью порождающей матрицы $\mathbf{G}$: множество χ является информационной совокуп-

ностью тогда и только тогда, когда $\text{rank } \mathbf{G}(\chi) = k$. Также можно использовать проверочную матрицу $\mathbf{H}$, в этом случае должно выполняться условие: $\text{rank } \mathbf{H}(\{1, \dots, n\} \setminus \chi) = r$, где дополнение $\{1, \dots, n\} \setminus \chi$ — множество чисел от 1 до n , не вошедших в χ .

Таким образом, если χ — информационная совокупность, матрица $\mathbf{G}(\chi)$ является невырожденной и к ней существует обратная (над полем $\text{GF}(2)$). Пусть $\mathbf{G}_\chi = \mathbf{G}^{-1}(\chi)\mathbf{G}$, тогда $\mathbf{G}_\chi$ также является базисом того же пространства, которое задается базисом $\mathbf{G}$, и, следовательно, $\mathbf{G}_\chi$ — еще одна порождающая матрица того же кода. На позициях χ матрица $\mathbf{G}_\chi$ содержит единичную подматрицу, следовательно, произведение $\mathbf{a}(\chi)\mathbf{G}_\chi$ задаст кодовое слово $\mathbf{a}$ с элементами $\mathbf{a}(\chi)$ на позициях χ . Декодирование на основе информационных совокупностей построено на поиске такого множества χ , которое не искажено ошибками, тогда кодовое слово $\mathbf{a}$ будет восстановлено правильно. Псевдокод алгоритма исправления независимых ошибок при декодировании по информационным совокупностям приведен на рис. 1.

Предполагается, что для использования алгоритма на рис. 1 необходимо сгенерировать множество информационных совокупностей X мощностью N , а также соответствующее множество порождающих матриц $\{\mathbf{G}_{\chi_1}, \dots, \mathbf{G}_{\chi_N}\}$. К сожалению, при исправлении независимых ошибок для обеспечения малых вероятностей ошибки декодирования число N должно быть экспоненциально велико.

Версия алгоритма, приведенная на рис. 1, не определяет непосредственно отсутствие ошибок на информационной совокупности, как в некоторых других вариациях, так как это требует знания значения t , которое в общем случае для случайных линейных кодов неизвестно, и вместо этого вычисляет функцию расстояния между векторами $d(\mathbf{a}, \mathbf{b})$. Из рассмотренных кодовых слов выбирается ближайшее к принятому слову. Для независимых битовых ошибок используется расстояние Хемминга. Данный алгоритм мож-

Вход: принятое из канала слово $\mathbf{b}$, множество $X = \{\chi_1, \dots, \chi_N\}$, множество $\{\mathbf{G}_{\chi_1}, \dots, \mathbf{G}_{\chi_N}\}$.

Выход: решение декодера о кодовом слове $\hat{\mathbf{a}}$.

1. $d_{\min} \leftarrow \infty$.

2. Для $i = 1$ до N :

2.1. $\mathbf{a} \leftarrow \mathbf{b}(\chi_i)\mathbf{G}_{\chi_i}$.

2.2. Если $d(\mathbf{a}, \mathbf{b}) < d_{\min}$:

2.2.1. $d_{\min} \leftarrow d(\mathbf{a}, \mathbf{b})$.

2.2.2. $\hat{\mathbf{a}} \leftarrow \mathbf{a}$.

3. Возвратить $\hat{\mathbf{a}}$.

■ **Рис. 1.** Декодирование по информационным совокупностям

■ **Fig. 1.** Information set decoding

но приспособить для исправления одиночных пакетов ошибок, однако необходимо учитывать специфику шума при построении множества X , а также при вычислении функции расстояния.

Для двух векторов $\mathbf{a}$ и $\mathbf{b}$ длиной n определим

$$d_b(\mathbf{a}, \mathbf{b}) = W_b(\mathbf{a} \oplus \mathbf{b}) = W_b(\mathbf{e}), \quad (1)$$

где $W_b(\mathbf{e})$ — функция, возвращающая длину пакета, который образует вектор $\mathbf{e}$. Тогда если $\mathbf{a}$ — передававшееся кодовое слово, $\mathbf{a}$ $\mathbf{b}$ — слово, принятое из канала, то $\mathbf{e} = \mathbf{a} \oplus \mathbf{b}$ — вектор ошибки. В этом случае алгоритм на рис. 1 будет принимать решение в пользу кодового слова, которому в канале связи соответствует пакет ошибок меньшей длины. Такой критерий декодирования ориентирован на обеспечение низких вероятностей ошибки в каналах, в которых более длинные пакеты ошибок менее вероятны, это согласуется с характеристиками многих реальных каналов связи.

При построении множества X и определении его объема N необходимо учитывать возможность нахождения информационной совокупности, не затронутой ошибками. В случае, когда ошибки образуют пакеты длиной b_0 (или менее) и они имеют возможные начала от 1 до $n - b_0 + 1$, нетрудно указать способ построения X , при котором гарантировано нахождение информационной совокупности, лежащей за пределами пакета. Для этого достаточно определить информационную совокупность, не включающую позиции пакета, тогда для всех возможных начал пакета $N = n - b_0 + 1$ и количество используемых информационных совокупностей линейно зависит от длины кода n .

Тем не менее возможно дальнейшее уменьшение величины N , что будет показано ниже.

Заметим, что функция расстояния, задаваемая (1), не является метрикой, как, например, расстояние Хемминга, и это не позволяет связать корректирующую способность b_0 со значениями длин пакетов $W_b(\mathbf{a})$, образуемых кодовыми словами, и утверждать о способности предложенного алгоритма реализовывать корректирующую способность кода при исправлении пакетов ошибок, что рассмотрим в следующем разделе.

Правило декодирования для исправления пакетов информационными совокупностями

При исправлении независимых ошибок с помощью поиска кодового слова, ближайшего к принятому, в метрике Хемминга декодирование будет гарантированно правильным, если число произошедших ошибок не превышает кор-

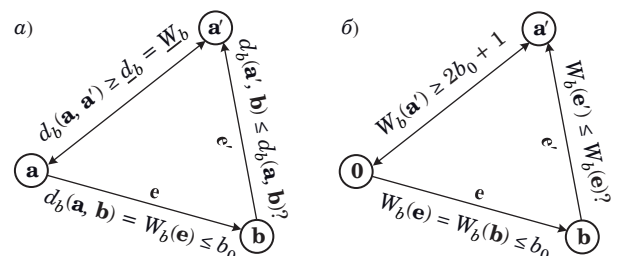
ректирующей способности t кода. Это используется и в алгоритме, приведенном на рис. 1: никакая информационная совокупность, искаженная не более чем t ошибками, не даст кодового слова, более близкого к принятому, чем правильное, и ошибка декодирования в данном алгоритме возможна лишь в случае, когда среди N информационных совокупностей не нашлось не затронутой ошибками.

Однако функция расстояния по длине пакета, определенная в (1), не является метрикой, в частности, для нее не выполняется неравенство треугольника, следовательно, нельзя напрямую связать корректирующую способность кода и декодирование по правилу (1). В данном разделе подробнее рассмотрим использование правила (1) в алгоритме на рис. 1.

В качестве иллюстрации приведем следующую ситуацию. Пусть C — (n, k) -код, способный исправлять любые векторы ошибки, представляющие собой пакеты длиной до b_0 . Это означает, что все такие векторы ошибки лежат в разных смежных классах. Пусть передается кодовое слово $\mathbf{a}$, и в канале произошла ошибка, описываемая вектором $\mathbf{e}$ и представляющая собой пакет длиной не более b_0 . Пусть из канала принято слово $\mathbf{b} = \mathbf{a} \oplus \mathbf{e}$. Рассмотрим возможность декодирования вектора $\mathbf{b}$ в неправильное кодовое слово $\mathbf{a}'$, при котором используется критерий минимума расстояния по правилу (1). Это означало бы, что декодер нашел вектор ошибки $\mathbf{e}' = \mathbf{a}' \oplus \mathbf{b}$ такой, что ошибочное кодовое слово $\mathbf{a}'$ ближе (или, возможно, не дальше) к принятому, чем правильное слово $\mathbf{a}$, или $d_b(\mathbf{a}', \mathbf{b}) \leq d_b(\mathbf{a}, \mathbf{b})$. Данная ситуация изображена на рис. 2, а.

Рассмотрим несколько простых соотношений для функции расстояния (1), многие из которых полностью эквивалентны таковым для метрики Хемминга. Для двоичных векторов $\mathbf{a}, \mathbf{b}, \mathbf{c}$ справедливо

$$d_b(\mathbf{a} \oplus \mathbf{c}, \mathbf{b} \oplus \mathbf{c}) = W_b(\mathbf{a} \oplus \mathbf{b}) = d_b(\mathbf{a}, \mathbf{b}).$$



■ **Рис. 2.** Оценка возможности ошибки при декодировании по критерию минимальной длины пакета в общем случае (а) и при передаче нулевого слова (б)

■ **Fig. 2.** Estimation of error possibility using minimal burst length decoding criteria in general case (a) and in case of zero codeword (b)

Обозначим через W_b минимальную длину пакета, которую образуют ненулевые кодовые слова кода C :

$$W_b = \min_{\mathbf{a} \in C, \mathbf{a} \neq \mathbf{0}} \{W_b(\mathbf{a})\}.$$

Обозначим через d_b минимальное попарное расстояние среди всех кодовых слов:

$$d_b = \min_{\substack{\mathbf{a}_i, \mathbf{a}_j \in C, \\ i \neq j}} \{d_b(\mathbf{a}_i, \mathbf{a}_j)\}.$$

Если $\mathbf{a}_i, \mathbf{a}_j$ — ближайшие кодовые слова, то

$$d_b = d_b(\mathbf{a}_i, \mathbf{a}_j) = W_b(\mathbf{a}_i \oplus \mathbf{a}_j) \geq W_b,$$

так как $\mathbf{a}_i \oplus \mathbf{a}_j$ также является кодовым словом. С другой стороны, для кодового слова $\mathbf{a}$ минимальной пакетной длины

$$W_b = W_b(\mathbf{a}) = d_b(\mathbf{0}, \mathbf{a}) \geq d_b,$$

следовательно: $W_b = d_b$.

Величина d_b имеет некоторую, хотя и не полную, аналогию с минимальным расстоянием d_0 кода (в метрике Хемминга), а W_b — с минимальным весом. Из теории кодирования известно, что, если код исправляет пакеты длиной b_0 и меньше, в коде нет ненулевых кодовых слов, образующих пакет длиной $2b_0$ или меньше, таким образом:

$$W_b \geq 2b_0 + 1. \quad (2)$$

Отличие от случая независимых ошибок вызвано тем, что расстояние (1) не является метрикой и состоит в том, что минимальное расстояние кода в метрике Хемминга и число независимых исправляемых ошибок связаны равенством $d_0 = 2t + 1$, тогда как в (2) присутствует неравенство. Таким образом, по величине b_0 нельзя точно определить W_b и наоборот.

Пусть $\mathbf{a}$ — вектор длиной n . Представим его как конкатенацию двух векторов: $\mathbf{a} = (\mathbf{a}_1 | \mathbf{a}_2)$, и образуем два вектора $\mathbf{e}_1 = (\mathbf{a}_1 | \mathbf{0})$ и $\mathbf{e}_2 = (\mathbf{0} | \mathbf{a}_2)$ длиной n каждый. Назовем множество $\{\mathbf{e}_1, \mathbf{e}_2\}$ декомпозицией вектора $\mathbf{a}$. Для дальнейшего рассмотрения докажем следующую лемму.

Лемма. Пусть $\mathbf{a}$ — кодовое слово. Тогда для любой его декомпозиции $\{\mathbf{e}_1, \mathbf{e}_2\}$ векторы $\mathbf{e}_1$ и $\mathbf{e}_2$ имеют одинаковый синдром.

Доказательство: Пусть $\mathbf{H}$ — проверочная матрица кода, тогда $\mathbf{aH}^T = \mathbf{0}$. Пусть $\mathbf{S}_1 = \mathbf{e}_1 \mathbf{H}^T$ и $\mathbf{S}_2 = \mathbf{e}_2 \mathbf{H}^T$, тогда $\mathbf{S}_1 \oplus \mathbf{S}_2 = (\mathbf{e}_1 \oplus \mathbf{e}_2) \mathbf{H}^T = \mathbf{aH}^T = \mathbf{0}$ или $\mathbf{S}_1 = \mathbf{S}_2$, что и требовалось доказать.

Из данной леммы прямо следует выражение (2). Действительно, если ненулевое кодовое слово $\mathbf{a}$ образует пакет длиной $2b_0$ или меньше, его де-

композиция, проведенная по средней позиции пакета, образует два пакета длиной, не превышающей b_0 , каждый, они будут иметь одинаковый синдром и находиться в одном смежном классе, что противоречит утверждению о том, что код исправляет пакеты длиной b_0 .

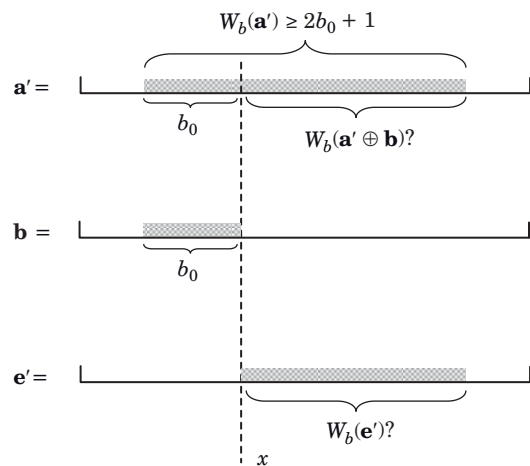
Подытоживая вышесказанное, прибавив (покомпонентно по модулю 2) ко всем векторам, изображенным на рис. 2, $\mathbf{a}$, вектор $\mathbf{a}$, мы получим более простой для рассмотрения случай (рис. 2, б), когда передается нулевое слово, а векторы $\mathbf{b}$ и $\mathbf{a}'$ заново переобозначены. При этом сохраняются все попарные расстояния между рассматриваемыми векторами, принятое слово $\mathbf{b}$ образует пакет длиной не более b_0 , а кодовое слово $\mathbf{a}'$ образует пакет длиной не менее $2b_0 + 1$. Сформулируем и докажем следующую теорему.

Теорема. Пусть код C исправляет все пакеты длиной до b_0 . При декодировании по минимуму расстояния с использованием функции расстояния (1) любой пакет ошибок длиной до b_0 будет корректно исправлен.

Доказательство: Чтобы доказать теорему, необходимо показать, что для ситуации, изображенной на рис. 2, б, невозможно выполнение неравенства $d_b(\mathbf{a}', \mathbf{b}) \leq d_b(\mathbf{a}, \mathbf{b})$ и, следовательно, невозможно выполнение $W_b(\mathbf{e}') \leq W_b(\mathbf{e})$.

Рассмотрим вектор $\mathbf{a}'$ и условия, необходимые для выполнения указанных неравенств. Вектор $\mathbf{a}'$ представляет собой одиночный пакет длиной не менее $2b_0 + 1$ (рис. 3).

Вектор $\mathbf{b}$ представляет собой пакет длиной не более b_0 , и нам необходимо определить, при каких условиях длина пакета вектора $\mathbf{e}' = \mathbf{a}' \oplus \mathbf{b}$ может оказаться меньше или равной длине пакета вектора $\mathbf{e}$. Другими словами, нужно определить возможность уменьшения пакета



■ **Рис. 3.** Уменьшение длины пакета при сложении векторов

■ **Fig. 3.** Decreasing of burst length from vectors addition

в векторе $\mathbf{a}'$ путем прибавления к нему вектора $\mathbf{b}$ (см. рис. 3).

Будем считать, что любой пакет начинается и заканчивается единицей, и нетрудно видеть, что максимальное уменьшение длины пакета может происходить, если пакет в векторе $\mathbf{b}$ полностью совпадает либо с левой частью пакета в векторе $\mathbf{a}'$ (что изображено на рис. 3), либо с правой, при этом длина пакета вектора $\mathbf{b}$ должна быть как можно большей, поэтому положим ее равной b_0 .

Рассмотрим сокращение пакета в векторе $\mathbf{a}'$ на b_0 позиций слева. Случай сокращения справа полностью аналогичен. Предположим, что в векторе $\mathbf{b}$ последняя позиция пакета равна $x - 1$, тогда пакет в векторе $\mathbf{e}'$ может начинаться с позиции x (это обозначено пунктирной линией на рис. 3). Неопределенность состоит в том, что элемент на позиции x в векторе $\mathbf{e}'$ может оказаться нулевым, как и последующие за ним, кроме последней позиции пакета, и длина пакета, таким образом, может оказаться сколь угодно маленькой, вплоть до единичной.

Воспользуемся доказанной леммой. Проведем декомпозицию вектора $\mathbf{a}'$ по позиции x , получив векторы $\mathbf{e}_1$ и $\mathbf{e}_2$, причем $\mathbf{e}_1 = \mathbf{e}$, $\mathbf{e}_2 = \mathbf{e}'$. Очевидно: $W_b(\mathbf{e}_1) = b_0$. По лемме, вектор $\mathbf{e}_2$ имеет тот же синдром, что и вектор $\mathbf{e}_1$, следовательно, находится с ним в одном смежном классе. Так как код исправляет пакеты длиной b_0 , в смежном классе может быть только один вектор длиной, не превышающей b_0 , следовательно: $W_b(\mathbf{e}_2) = W_b(\mathbf{e}') \geq b_0 + 1$ и $W_b(\mathbf{e}') > W_b(\mathbf{e})$, что и требовалось доказать.

Следствие. Если длина пакета ошибок не превышает корректирующей способности кода b_0 , декодирование по информационным совокупностям, использующее функцию расстояния (1), гарантированно исправит пакет ошибок, если в множестве X , заданном алгоритму на рис. 1, окажется хотя бы одна информационная совокупность, не содержащая позиций пакета. Это можно обеспечить, включив в X хотя бы одну такую информационную совокупность для каждой из $n - b_0 + 1$ начальных позиций пакета.

Анализ методик построения информационных совокупностей

В предыдущем разделе мы доказали, что алгоритм декодирования одиночных пакетов по информационным совокупностям, использующий правило (1), гарантированно обеспечивает декодирование в пределах корректирующей способности кода b_0 . При этом множество информационных совокупностей X должно содержать $n - b_0 + 1 = O(n)$ информационных совокупностей. Однако, хотя в общем случае нахождение информационной совокупности на случайной

позиции имеет достаточно большую вероятность, при ограничении выбираемых позиций и при снижении числа информационных совокупностей необходимо оценивать вероятность построения того или иного множества с учетом ограничений и конструкции кода.

В этом разделе рассмотрим подходы к дальнейшему уменьшению количества информационных совокупностей и оценке вероятности нахождения множества X для случайных линейных кодов.

Обратим внимание, что информационная совокупность, образованная, к примеру, числами $\{1, \dots, k\}$, будет свободной от ошибок для любого пакета, начинающегося с позиции $k + 1$. Таким образом, останется найти только информационные совокупности для пакетов, начало которых приходится на первые k позиций. Такой подход может позволить снизить количество информационных совокупностей до некоторого значения $T \leq N = n - b_0 + 1$.

Конечно, нет никаких гарантий, что k последовательно идущих чисел образуют информационную совокупность, поэтому введем параметр $\Delta \in [0, n - k - b_0]$ и будем искать информационные совокупности в интервале длиной $k + \Delta$. Значение $k + \Delta$ назовем диаметром информационной совокупности. Ограниченные таким образом совокупности назовем плотными [14]. Нетрудно показать, что при $\Delta = n - k - b_0$ интервалы поиска информационных совокупностей будут ограничены только длинами пакетов, что совпадает с исходной процедурой поиска.

Вероятность нахождения информационной совокупности на интервале $k + \Delta$ соответствует вероятности нахождения невырожденной матрицы $\mathbf{M}$ размера $k \times k$. Известно [21], что для случайной $(k \times k + \Delta)$ -матрицы эта вероятность может быть примерно оценена как

$$\Pr(\text{rank}(\mathbf{M}_{k,k+\Delta}) = k) \approx Q_\Delta = \prod_{j=\Delta+1}^{\infty} \left(1 - \frac{1}{2^j}\right) \quad (3)$$

для всех натуральных k , имеющих практический интерес (больше нескольких единиц).

В частности, при $\Delta = 0$ вероятность выбора случайной невырожденной двоичной $(k \times k)$ -матрицы составляет примерно 0,28 вне зависимости от значения k . Предположим, что нахождение информационных совокупностей представляет собой последовательность независимых событий, тогда из (3) следует, что вероятность нахождения множества из T информационных совокупностей будет иметь следующий вид:

$$Q_{\Delta T} = Q_\Delta^T = \left(\prod_{j=\Delta+1}^{\infty} \left(1 - \frac{1}{2^j}\right) \right)^T. \quad (4)$$

Сопоставим некоторой плотной информационной совокупности множество начал пакетов, для которых пакет не затрагивает данную информационную совокупность. Назовем позиции из такого множества защищенными (данной информационной совокупностью). Тогда процедура построения множества X информационных совокупностей должна продолжаться до тех пор, пока защищенными не станут все возможные начала пакетов, т. е. $\{1, \dots, n - b_0 + 1\}$.

Рассмотрим методику построения, при которой начала плотных информационных совокупностей для случайного линейного кода выбираются циклически с шагом $k + \Delta + b_0$. Если на очередном интервале длиной $k + \Delta$ информационная совокупность не найдена, будем считать, что происходит отказ процедуры поиска. На практике в этом случае можно либо сгенерировать другой случайный код, либо расширить значение Δ .

Число плотных информационных совокупностей, построенных по такой методике при заданных Δ и b_0 , можно оценить с помощью выражения

$$T \geq \left\lceil \frac{n - b_0 + 1}{n - b_0 + 1 - (k + \Delta)} \right\rceil. \quad (5)$$

Если считать, что используются коды, корректирующая способность которых соответствует границе Рейгера ($b_0 = (n - k)/2$), то выражение (5) может быть записано как

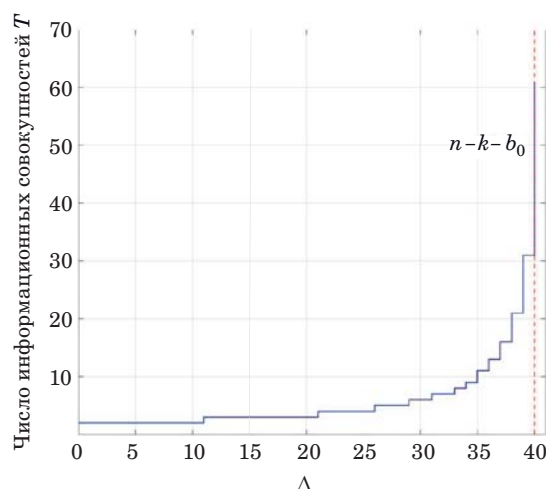
$$T \geq \left\lceil \frac{n + nR + 2}{n + nR + 2 - 2(nR + \Delta)} \right\rceil, \quad (6)$$

где $R = k/n$ — скорость кода.

Как нетрудно убедиться, при фиксированных скорости кода и параметре Δ число T плотных информационных совокупностей оценивается как $O(1)$ от длины кода n . При фиксированных n и Δ величина T растет с ростом R . Наконец, используя выражения (4) и (6), рассмотрим зависимость величины T от значения Δ при фиксированных n и R , например, для кода со скоростью $R = 0,2$, длиной $n = 100$, $b_0 = 40$ (рис. 4).

График зависимости имеет ступенчатый вид, и можно наблюдать достаточно большие диапазоны значений Δ , которым соответствует одинаковое число информационных совокупностей T . Рассмотрим значения функции Q_{Δ}^T (4) — вероятности нахождения множества плотных информационных совокупностей — при фиксированном T . Функция Q_{Δ}^T из выражения (4) — возрастающая, рассмотрим ее значения, например при $T = 2$, в зависимости от значений Δ (таблица).

Как можно видеть, при малых Δ вероятность построения множества X плотных информаци-



■ Рис. 4. Зависимость числа плотных информационных совокупностей T от значения Δ

■ Fig. 4. Dependence of the number T of dense information sets from the Δ value

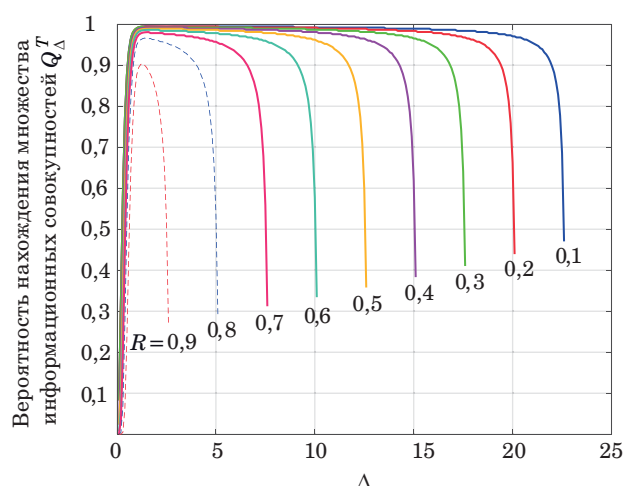
■ Значения функции Q_{Δ}^T для $T = 2$ и кода: $n = 100$, $R = 0,2$

■ Function values Q_{Δ}^T for $T = 2$ and code with parameters: $n = 100$, $R = 0.2$

Δ	Q_{Δ}^2	Δ	Q_{Δ}^2
0	0,08	6	0,96
1	0,33	7	0,98
2	0,6	8	0,99
3	0,87	9	0,99
4	0,88	10	0,99
5	0,93		

онных совокупностей не очень велика, однако с ростом Δ она приближается к единице.

Следовательно, при построении множества плотных информационных совокупностей можно выбирать то значение Δ , при котором вероятность Q_{Δ}^T максимальна, т. е. максимальное для заданного числа T . Рассмотрим теперь зависимость вероятности нахождения множества информационных совокупностей от значения Δ для разных скоростей кодов с длиной $n = 100$ (рис. 5). Приведенные зависимости не являются монотонными, и существует некоторое оптимальное значение Δ , максимизирующее значение Q_{Δ}^T (его довольно нетрудно найти по приведенным формулам). Заметим, что с увеличением скорости кода это оптимальное значение заметно снижается, это вызвано тем, что для высокоскоростных кодов значение T достаточно велико.



■ **Рис. 5.** Зависимость вероятности нахождения множества информационных совокупностей Q_A^T от значения Δ при разных скоростях R кода

■ **Fig. 5.** Dependence of dense information set probability Q_A^T from values of Δ for different code rates R

Заключение

В статье рассмотрена задача исправления одиночных пакетов случайными линейными кодами с помощью подхода на основе информационных совокупностей. Предложен вариант декодирования по функции расстояния, соответствующей длине пакета. Доказано, что декодирование по минимуму длины пакета гарантирует исправление всех пакетов в пределах корректирующей способности кода. На основании этого доказано, что при использовании множества

информационных совокупностей, покрывающих все начала пакетов, декодирование по информационным совокупностям и минимуму длины пакета реализует корректирующую способность кода. При этом число необходимых информационных совокупностей оценивается как $O(n)$.

Рассмотрена методика построения множества плотных информационных совокупностей, позволяющая уменьшить их число до $O(1)$. Проведен анализ вероятности построения множества плотных информационных совокупностей в зависимости от их диаметра и предложен способ выбора диаметра, на основе которого сформулирована оптимизационная задача, позволяющая максимизировать вероятность построения множества информационных совокупностей.

Все вероятностные расчеты проведены в предположении, что рассматривается случайный линейный двоичный код. Возможным направлением дальнейших исследований является распространение подходов, рассмотренных в статье, на конкретные классы помехоустойчивых кодов, которые обладают достаточно хорошей корректирующей способностью для исправления одиночных пакетов ошибок.

Финансовая поддержка

Работа выполнена при финансовой поддержке Российского научного фонда, грант № 22-19-00305 «Пространственно-временные стохастические модели беспроводных сетей с большим числом абонентов».

Литература

1. Yang M., Pan Z., Djordjevic I. B. FPGA-based burst-error performance analysis and optimization of regular and irregular SD-LDPC codes for 50G-PON and beyond. *Opt. Express*, 2023, vol. 31, no. 6, pp. 10936–10946. doi:10.1364/OE.477546
2. Song L., Huang Q., Wang Z. Construction of multiple-burst-correction codes in transform domain and its relation to LDPC codes. *IEEE Trans. Commun.*, 2020, vol. 68, no. 1, pp. 40–54. doi:10.1109/TCOMM.2019.2948341
3. Vafi S. Cyclic low density parity check codes with the optimum burst error correcting capability. *IEEE Access*, 2020, vol. 8, pp. 192065–192072. doi:10.1109/ACCESS.2020.3032837
4. Xiao X., Vasic B., Lin S., Li J., Abdel-Ghaffar K. Quasi-cyclic LDPC codes with parity-check matrices of column weight two or more for correcting phased bursts of erasures. *IEEE Trans. Commun.*, 2021, vol. 69, no. 5, pp. 2812–2823. doi:10.1109/TCOMM.2021.3059001
5. Aharoni Z., Huleihel B., Pfister H. D., Permuter H. H. Data-driven polar codes for unknown channels with and without memory. *2023 IEEE International Symposium on Information Theory (ISIT)*, Taipei, IEEE, 2023, pp. 1890–1895. doi:10.1109/ISIT54713.2023.10206663
6. Sasoglu E., Tal I. Polar coding for processes with memory. *IEEE Trans. Inform. Theory*, 2019, vol. 65, no. 4, pp. 1994–2003. doi:10.1109/TIT.2018.2885797
7. Vajha M., Ramkumar V., Jhamtani M., Kumar P. V. On the performance analysis of streaming codes over the Gilbert – Elliott channel. *2021 IEEE Information Theory Workshop (ITW)*, Kanazawa, Japan, 2021, pp. 1–6. doi:10.1109/ITW48936.2021.9611448
8. Eckford A., Kschischang F., Pasupathy S. Analysis of low-density parity-check codes for the Gilbert – Elliott channel. *IEEE Transactions on Information Theory*, 2005, vol. 51, no. 11, pp. 3872–3889. doi:10.1109/TIT.2005.856934
9. Li L., Lv J., Li Y., Dai X., Wang X. Burst error identification method for LDPC coded systems. *IEEE Com-*

- mun. Lett.*, 2024, pp. 1–5. doi:10.1109/LCOMM.2024.3391826
10. Fang Y., Chen J. Decoding polar codes for a generalized Gilbert – Elliott channel with unknown parameter. *IEEE Trans. Commun.*, 2021, vol. 69, no. 10, pp. 6455–6468. doi:10.1109/TCOMM.2021.3095195
 11. Трофимов А. Н., Таубин Ф. А. Улучшенная граница вероятности ошибки при оптимальном приеме в канале с межсимвольной интерференцией. *Информационно-управляющие системы*, 2023, № 5, с. 33–42. doi:10.31799/1684-8853-2023-5-33-42, EDN: MDHOXU
 12. Ovchinnikov A. A., Veresova A. M., Fominykh A. A. Decoding of linear codes for single error bursts correction based on the determination of certain events. *Информационно-управляющие системы*, 2022, № 6, с. 41–52. doi:10.31799/1684-8853-2022-6-41-52, EDN: UWXZHN
 13. Veresova A., Fominykh A., Ovchinnikov A. About usage of metrics in decoding of LDPC codes in two-state channels with memory. *2021 XVII International Symposium Problems of Redundancy in Information and Control Systems (REDUNDANCY)*, Moscow, IEEE, 2021, pp. 143–148. doi:10.1109/REDUNDANCY52534.2021.9606474
 14. Исаева М. Н., Овчинников А. А. Исправление одиночных пакетов ошибок за пределами корректирующей способности кода с использованием информационных совокупностей. *Научно-технический вестник информационных технологий, механики и оптики*, 2024, вып. 24, № 1, с. 70–80. doi:10.17586/2226-1494-2024-24-1-70-80
 15. Sung W., Coffey J. T. Information set decoding complexity for linear codes in bursty channels with side information. *Proceedings of 1995 IEEE International Symposium on Information Theory*, Whistler, BC, Canada, 1995, pp. 53. doi:10.1109/ISIT.1995.531155
 16. Moon T. K. *Error Correction Coding: Mathematical Methods and Algorithms*. Second ed. Hoboken, NJ, Wiley, 2021. 992 p.
 17. Gazi O. *Forward Error Correction Via Channel Coding*. Cham, Springer, 2020. 319 p.
 18. Lin S., Li J. *Fundamentals of Classical and Modern Error-Correcting Codes*. Cambridge, Cambridge University Press, 2022. 840 p. doi:10.1017/9781009067928
 19. Veresova A. M., Isaeva M. N., Ovchinnikov A. A. Estimation of independent errors and bursts correction capability of linear codes. *2024 Conference of Young Researchers in Electrical and Electronic Engineering (ElCon)*, Saint-Petersburg, IEEE, 2024, pp. 23–27. doi:10.1109/ElCon61730.2024.10468456
 20. Исаева М. Н. Поиск информационных совокупностей при исправлении пакетов ошибок квазициклическими кодами. *Т-Сотт: Телекоммуникации и транспорт*, 2023, т. 17, № 7, с. 4–12. doi:10.36724/2072-8735-2023-17-7-4-12
 21. Berlekamp E. R. The technology of error correcting codes. *Proc. IEEE*, 1980, vol. 68, pp. 564–593.

UDC 519.72

doi:10.31799/1684-8853-2025-2-68-77

EDN: MAHCAY

Decoding of single error bursts using minimal burst length criteria and information setsM. N. Isaeva^a, Senior Lecturer, orcid.org/0009-0007-6228-0617, imn@guap.ru^aSaint-Petersburg State University of Aerospace Instrumentation, 67, B. Morskaya St., 190000, Saint-Petersburg, Russian Federation

Introduction: To improve the characteristics of data transmission and storage systems the specificity of noise should be taken into account, which tends to form the error bursts in many practical situations. Information set decoding for independent errors is the best decoding approach for random linear codes, but has exponential complexity. **Purpose:** To analyze the error-correcting capability and computational complexity of algorithms based on information sets in order to correct single error bursts. **Results:** We suggest the decoding criteria of minimal burst length for single bursts correction using information sets. We prove the correction of all error bursts within burst-correction capability of the code using this criterion. The task of decreasing the number of information sets required for correct decoding is considered, for this task the methodology of constructing the information sets with limited diameter, which are called dense sets, is investigated. We perform the analysis showing the possibility of decreasing the number of information sets from $O(n)$ to $O(1)$. We estimate the probability of information sets construction and formulate the optimization problem for maximizing this probability. **Practical relevance:** The developed decoder of single error bursts basing on minimal burst length criteria, as well as methodology of construction the collection of dense information sets with cardinality $O(1)$ guarantee the realization of burst-correction capability of random linear codes. **Discussion:** The results obtained show that the problem of single error bursts correction with random linear codes is polynomial. At the same time, the proposed methods make it possible to get a high probability construction of information sets for codes with a relatively low rate, and these probabilities can significantly change due to the use of more practical classes of codes, which may be considered as future research directions.

Keywords – information set decoding, Reiger bound, error burst correction, dense information set, decoding criteria.

For citation: Isaeva M. N. Decoding of single error bursts using minimal burst length criteria and information sets. *Informatsionno-upravliaiushchie sistemy* [Information and Control Systems], 2025, no. 2, pp. 68–77 (In Russian). doi:10.31799/1684-8853-2025-2-68-77, EDN: MAHCAY

Financial support

The paper was prepared with the financial support of the Russian Science Foundation, project No. 22-19-00305 “Spatial-temporal stochastic models of wireless networks with a large number of users”.

References

1. Yang M., Pan Z., Djordjevic I. B. FPGA-based burst-error performance analysis and optimization of regular and irregular SD-LDPC codes for 50G-PON and beyond. *Opt. Express*, 2023, vol. 31, no. 6, pp. 10936–10946. doi:10.1364/OE.477546
2. Song L., Huang Q., Wang Z. Construction of multiple-burst-correction codes in transform domain and its relation to LDPC codes. *IEEE Trans. Commun.*, 2020, vol. 68, no. 1, pp. 40–54. doi:10.1109/TCOMM.2019.2948341
3. Vafi S. Cyclic low density parity check codes with the optimum burst error correcting capability. *IEEE Access*, 2020, vol. 8, pp. 192065–192072. doi:10.1109/ACCESS.2020.3032837
4. Xiao X., Vasic B., Lin S., Li J., Abdel-Ghaffar K. Quasi-cyclic LDPC codes with parity-check matrices of column weight two or more for correcting phased bursts of erasures. *IEEE Trans. Commun.*, 2021, vol. 69, no. 5, pp. 2812–2823. doi:10.1109/TCOMM.2021.3059001
5. Aharoni Z., Huleihel B., Pfister H. D., Permuter H. H. Data-driven polar codes for unknown channels with and without memory. *2023 IEEE International Symposium on Information Theory (ISIT)*, Taipei, IEEE, 2023, pp. 1890–1895. doi:10.1109/ISIT54713.2023.10206663
6. Sasoglu E., Tal I. Polar coding for processes with memory. *IEEE Trans. Inform. Theory*, 2019, vol. 65, no. 4, pp. 1994–2003. doi:10.1109/TIT.2018.2885797
7. Vajha M., Ramkumar V., Jhamtani M., Kumar P. V. On the performance analysis of streaming codes over the Gilbert – Elliott channel. *2021 IEEE Information Theory Workshop (ITW)*, Kanazawa, Japan, 2021, pp. 1–6. doi:10.1109/ITW48936.2021.9611448
8. Eckford A., Kschischang F., Pasupathy S. Analysis of low-density parity-check codes for the Gilbert – Elliott channel. *IEEE Transactions on Information Theory*, 2005, vol. 51, no. 11, pp. 3872–3889. doi:10.1109/TIT.2005.856934
9. Li L., Lv J., Li Y., Dai X., Wang X. Burst error identification method for LDPC coded systems. *IEEE Commun. Lett.*, 2024, pp. 1–5. doi:10.1109/LCOMM.2024.3391826
10. Fang Y., Chen J. Decoding polar codes for a generalized Gilbert – Elliott channel with unknown parameter. *IEEE Trans. Commun.*, 2021, vol. 69, no. 10, pp. 6455–6468. doi:10.1109/TCOMM.2021.3095195
11. Trofimov A. N., Taubin F. A. Improved bound on optimal reception error probability for an intersymbol interference channel. *Informatsionno-upravliaiushchie sistemy* [Information and Control Systems], 2023, no. 5, pp. 33–42 (In Russian). doi:10.31799/1684-8853-2023-5-33-42, EDN: MDHOU
12. Ovchinnikov A. A., Veresova A. M., Fominykh A. A. Decoding of linear codes for single error bursts correction based on the determination of certain events. *Informatsionno-upravliaiushchie sistemy* [Information and Control Systems], 2022, no. 6, pp. 41–52. doi:10.31799/1684-8853-2022-6-41-52, EDN: UWXXZN
13. Veresova A., Fominykh A., Ovchinnikov A. About usage of metrics in decoding of LDPC codes in two-state channels with memory. *2021 XVII International Symposium Problems of Redundancy in Information and Control Systems (REDUNDANCY)*, Moscow, IEEE, 2021, pp. 143–148. doi:10.1109/REDUNDANCY52534.2021.9606474
14. Isaeva M. N., Ovchinnikov A. A. Correction of single error bursts beyond the code correction capability using information sets. *Scientific and Technical Journal of Information Technologies, Mechanics and Optics*, 2024, vol. 24, no. 1, pp. 70–80 (In Russian). doi:10.17586/2226-1494-2024-24-1-70-80
15. Sung W., Coffey J. T. Information set decoding complexity for linear codes in bursty channels with side information. *Proceedings of 1995 IEEE International Symposium on Information Theory*, Whistler, BC, Canada, 1995, pp. 53. doi:10.1109/ISIT.1995.531155
16. Moon T. K. *Error Correction Coding: Mathematical Methods and Algorithms*. Second ed. Hoboken, NJ, Wiley, 2021. 992 p.
17. Gazi O. *Forward Error Correction Via Channel Coding*. Cham, Springer, 2020. 319 p.
18. Lin S., Li J. *Fundamentals of Classical and Modern Error-Correcting Codes*. Cambridge, Cambridge University Press, 2022. 840 p. doi:10.1017/9781009067928
19. Veresova A. M., Isaeva M. N., Ovchinnikov A. A. Estimation of independent errors and bursts correction capability of linear codes. *2024 Conference of Young Researchers in Electrical and Electronic Engineering (ElCon)*, Saint-Petersburg, IEEE, 2024, pp. 23–27. doi:10.1109/ElCon61730.2024.10468456
20. Isaeva M. N. Finding information sets when correcting error bursts with quasi-cyclic codes. *T-Comm*, vol. 17, no. 7, pp. 4–12 (In Russian). doi:10.36724/2072-8735-2023-17-7-4-12
21. Berlekamp E. R. The technology of error correcting codes. *Proc. IEEE*, 1980, vol. 68, pp. 564–593.