

(m, n) –Closed and Quasi (m, n) –Closed Ideals

Erlangga Adinugroho Rohi^{1*}, Sri Wahyuni²

^{1,2}Department of Mathematics, Universitas Gadjah Mada, Indonesia

¹erlanggaadinugrohorohi1999@mail.ugm.ac.id, ²swahyuni@ugm.ac.id

Abstract. Throughout this paper, all rings considered are commutative rings R with identity 1_R . Let m and n be natural numbers such that $1 \leq n < m$. A proper ideal I of R is called an (m, n) –closed ideal if for every $x \in R$ with $x^m \in I$ implies $x^n \in I$. An (m, n) –closed ideal generalizes semi n –absorbing ideal and, hence, also generalizes semiprime ideal. A proper ideal I of R is called a quasi (m, n) –closed ideal if for every $x \in R$ with $x^m \in I$ implies $x^n \in I$ or $x^{m-n} \in I$. Therefore, a quasi (m, n) –closed generalizes an (m, n) –closed ideal. Research related to these ideals is referred to Anderson and Badawi (2017) and Khashan and Celikel (2024). In this paper, the authors presented several new properties related to these ideals that are not discussed in the two main references.

Keywords: (m, n) –closed ideal, quasi (m, n) –closed ideal, semiprime ideal, commutative ring with identity.

1. INTRODUCTION

The idea of ideal comes from the last Fermat theorem. Its proof became an open problem until Andrew Wiles found it in 1994. Before Andrew Wiles, Ernst Kummer (1810-1893) tried to solve it. During his work, he introduced a concept of "ideal number". Motivated from his idea, Richard Dedekind (1831-1916) invented ideal concept in ring theory [1].

Prime ideal is one of many concepts in ring theory. It is motivated from the prime concept in number theory. Generalization of prime ideal has been found by many mathematician. For example, Anderson and Badawi [2] define an n –absorbing ideal as a generalization of prime ideal and conclude that a prime ideal is just 1–absorbing ideal. This ideal has been used by Choi [3] to prove Anderson and Badawi's conjectures in locally divided commutative rings. On the other

*Corresponding author

2020 Mathematics Subject Classification: 13A15, 13F20, 13G99, 13H99

Received: 12-05-2025, accepted: 08-07-2025.

hand, Moghimi and Naghani [4] define the n -Krull dimension of commutative ring R as a supremum of the lengths of chains of n -absorbing ideals of R .

There is another kind of ideal in ring theory. It is called a semiprime or radical ideal. Generally, a prime ideal is obvious a semiprime ideal. However, its converse need not to be true. Not only prime ideal, Anderson and Badawi [?] define a semi n -absorbing ideal as a generalization of semiprime ideal. They conclude that a semiprime ideal is just a semi 1-absorbing ideal. They also prove that every n -absorbing ideal is semi n -absorbing, but its converse need not to be true in general.

Anderson and Badawi [5] also generalize semi n -absorbing ideal. They define it as (m, n) -closed ideal for some natural numbers m and n . Therefore, an (m, n) -closed ideal also generalizes semiprime ideal. As a result of their research, one of its basic properties is every semiprime ideals are (m, n) -closed ideal for every natural numbers m and n .

Ideal generated by 2^{m-2} , i.e. $\langle 2^{m-2} \rangle$, is not $(m, 2)$ -closed ideal in $\mathbb{Z}$ for every natural numbers $m \geq 5$. Khashan and Celikel [6] define a generalization of (m, n) -closed ideal as quasi (m, n) -closed ideal. By defining this ideal, ideal $\langle 2^{m-2} \rangle$ is a quasi $(m, 2)$ -closed ideal in $\mathbb{Z}$ for every natural numbers $m \geq 5$. In addition, they also stated a sufficient condition for quasi (m, n) -closed ideal being (m, n) -closed ideal.

Many properties of (m, n) -closed ideal and quasi (m, n) -closed has been produced by [5] and [6] respectively. On the other hand, the applications of (m, n) -closed to several kind of rings also have been found by several papers. Issoual *et al.* [7] have investigated (m, n) -closed ideal related to the amalgamated ring $A \bowtie^f J$ for (A, B) be a pair of rings, $f : A \rightarrow B$ be a ring homomorphism and J be an ideal of B . Badawi *et al.* [8] also have explored (m, n) -closed ideal related to trivial ring extension $R = A(+)M$ for A is a commutative ring and M is an A -module. In this paper, the authors give several new properties related to these ideals that are not discussed in [5] and [6].

2. SOME CONCEPTS ABOUT (m, n) -CLOSED AND QUASI (m, n) -CLOSED IDEALS

In this section, we briefly give some explanations about (m, n) -closed and quasi (m, n) -closed ideal. Throughout this paper, all rings considered are commutative rings R with identity 1_R and all ring homomorphisms preserve the identity.

Let R be a commutative ring with identity. An ideal is proper if $I \neq R$. For some proper ideal I of R , the radical of I , denoted by $\sqrt{I}$, is defined by $\{x \in R \mid x^n \in I \text{ for some } n \in \mathbb{N}\}$. An element of R is called a nilpotent element if there exists a natural number n such that $x^n = 0_R$. The set of all nilpotent elements is called nilradical of R , denoted by $\text{nil}(R)$. Clearly, we have $\sqrt{\{0_R\}} = \text{nil}(R)$.

A proper ideal I of R is called a prime ideal if whenever $xy \in I$ for every $x, y \in R$, implies $x \in I$ or $y \in I$ [9]. Ideal generated by a prime number p of $\mathbb{Z}$,

i.e. $\langle p \rangle$, is an example of prime ideal. A semiprime ideal is a proper ideal with the property that whenever $x^2 \in I$ for every $x \in R$, implies $x \in I$ [9]. Every prime ideal is semiprime ideal, but its converse need not to be true. Ideal generated by $\langle 6 \rangle$ of $\mathbb{Z}$, i.e. $\langle 6 \rangle$, is a semiprime ideal of $\mathbb{Z}$ but it is not a prime ideal of $\mathbb{Z}$. Anderson and Badawi [2] generalize a prime ideal to n -absorbing ideal.

Definition 2.1. [2] *Let n be a natural number. A proper ideal I of R is called an n -absorbing ideal of R if whenever $x_1 x_2 \dots x_n x_{n+1} \in I$ for every $x_1, x_2, \dots, x_n, x_{n+1} \in R$, then there are n of x_i 's whose product is in I .*

In particular, an n -absorbing ideal is expanded from 2-absorbing ideal concept that has been explored by Badawi [10] and Payrovi and Babaei [?]. Same with prime ideal, Anderson and Badawi [5] generalize a semiprime ideal to semi n -absorbing ideal.

Definition 2.2. [5] *Let n be a natural number. A proper ideal I of R is called a semi n -absorbing ideal if whenever $x^{n+1} \in I$ for every $x \in R$, implies $x^n \in I$.*

Thus, an n -absorbing ideal is a semi n -absorbing ideal. It follows from definitions, a prime ideal and a semiprime ideal are a 1-absorbing ideal and a semi n -absorbing ideal, respectively.

Anderson and Badawi [5] also generalize a semi n -absorbing to an (m, n) -closed ideal.

Definition 2.3. [5] *Let m and n be natural numbers such that $1 \leq n < m$. A proper ideal I of R is called an (m, n) -closed ideal if whenever $x^m \in I$ for every $x \in R$, implies $x^n \in I$.*

Note that Mostafanasab and Darani [11] define a proper ideal I of R to be a semi (m, n) -absorbing ideal if I is an (m, n) -closed ideal. Here is an example of an (m, n) -closed ideal.

Example 2.4. *Let $R = \mathbb{Z}[x, y]$ and $I = \langle x^2, 2xy, y^2 \rangle$. We will prove that I is an $(m, 2)$ -closed ideal of R for any natural numbers $m \geq 3$. Suppose $f(x, y) \in R$ such that $(f(x, y))^m \in I$. We have $f(x, y) \in \sqrt{I} = \langle x, y \rangle$. Then, $f(x, y)$ can be written as $f(x, y) = g(x, y)x + h(x, y)y$ for some $g(x, y), h(x, y) \in R$. By squaring both sides, we get $(f(x, y))^2 = (g(x, y))^2 x^2 + 2xyg(x, y)h(x, y) + (h(x, y))^2 y^2 \in I$.*

By induction, we can prove that a semiprime ideal is an (m, n) -closed ideal for every natural numbers m and n . Note that an (m, n) -closed ideal is also an (m', n') -closed ideal for every natural numbers $m' \leq m$ and $n' \geq n$. The complete properties of an (m, n) -closed ideal can be found on [5].

The concept of (m, n) -closed ideal is generalized by Khashan and Celikel [6]. They define it as a quasi (m, n) -closed ideal.

Definition 2.5. [6] *Let m and n be natural numbers such that $1 \leq n < m$. A proper ideal I of R is called a quasi (m, n) -closed ideal if whenever $x^m \in I$ for every $x \in R$, implies $x^n \in I$ or $x^{m-n} \in I$.*

Here is an example of a quasi (m, n) -closed ideal.

Example 2.6. Let $R = \mathbb{Z}$. We will show that $I = \langle 2^{m-2} \rangle$ is a quasi $(m, 2)$ -closed ideal of R for every natural numbers $m \geq 5$. Let $x \in R$ such that $x^m \in I$. By the definition of I , we have $2^{m-2} | x^m$. This implies that $2 | x$. Consequently, we get $x^{m-2} \in I$. However, I is not an $(m, 2)$ -closed ideal since $2^m \in I$ but $2^2 \notin I$.

It is clear that a proper ideal I of R is a quasi (m, n) -closed ideal if and only if I is an (m, n) -closed or an $(m, m-n)$ -closed ideal. Furthermore, a proper ideal I of R is a quasi (m, n) -closed ideal if and only if I is a quasi $(m, m-n)$ -closed ideal. The comprehensive properties about quasi (m, n) -closed ideal can be found on [6].

3. RESULT AND DISCUSSION

In this section, we deliver several new properties related to (m, n) -closed ideals and quasi (m, n) -closed ring that are not discussed in [5] and [6].

3.1. (m, n) -closed Ideal. Let m and n be natural numbers such that $1 \leq n < m$. Corollary 2.4 in [5] shows that if $I_1, I_2, \dots, I_k$ are (m, n) -closed ideals of a commutative ring R with identity 1_R , then the intersection $I_1 \cap I_2 \cap \dots \cap I_k$ is also an (m, n) -closed ideal of R . Fortunately, this fact can be also extended to the collection of (m, n) -closed ideals of R .

Theorem 3.1. Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $\{I_\alpha | \alpha \in \Lambda\}$, where $\emptyset \neq \Lambda$ denotes an indexing set, is a collection of (m, n) -closed ideals of R , then $\cap_{\alpha \in \Lambda} I_\alpha$ is an (m, n) -closed ideal of R .

Proof. Let $x \in R$ such that $x^m \in \cap_{\alpha \in \Lambda} I_\alpha$, then $x^m \in I_\alpha$ for every $\alpha \in \Lambda$. Since I_α is an (m, n) -closed ideal for every $\alpha \in \Lambda$, we have $x^n \in \cap_{\alpha \in \Lambda} I_\alpha$. $\square$

In commutative ring R , the set $\text{nil}(R)$ is an ideal of R . Its proof can be found on [12]. For an ideal I of R , the set $\sqrt{I}$ is also an ideal of R [13]. Now, we can prove that $\text{nil}(R)$ and $\sqrt{I}$ are (m, n) -closed ideals for every natural numbers m and n such that $1 \leq n < m$.

Theorem 3.2. If R is a commutative ring with identity 1_R , then $\text{nil}(R)$ and $\sqrt{I}$, for an ideal I of R , are (m, n) -closed ideals of R for every natural numbers m and n .

Proof. It is enough to show that $\text{nil}(R)$ and $\sqrt{I}$ are semiprime ideals of R . Let $x \in R$ such that $x^2 \in \sqrt{I}$. Based on the definition of $\sqrt{I}$, there is a natural number k such that $x^{2k} = (x^2)^k \in I$. Since $2k$ is also a natural number, we get $x \in \sqrt{I}$. Now, let $y \in R$ such that $y^2 \in \text{nil}(R)$. By definition, there is a natural number j such that $y^{2j} = (y^2)^j = 0_R$. Moreover, we get $y \in \text{nil}(R)$. $\square$

It is easy to verify that if I is a proper ideal of a commutative ring R with identity 1_R , then $I[x]$ is a proper ideal of a polynomial ring $R[x]$. It is also obvious that $I = I[x] \cap R$. Hence, we get a fact below. Definitely, this fact is a consequence of the first statement of Corollary 2.11 on [5].

Corollary 3.3. *Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $I[x]$ is an (m, n) -closed ideal of $R[x]$, then I is an (m, n) -closed ideal of R .*

Proof. Note that R can be embedded to $R[x]$, so we have $R \subseteq R[x]$. Applying the first statement of Corollary 2.19 on [5] and the fact $I = I[x] \cap R$, we have I is an (m, n) -closed ideal of R . $\square$

In order to prove our next results, we present the second statement of Corollary 2.11 on [5].

Corollary 3.4. [5] *Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $I \subseteq J$ are proper ideals of R , then J/I is an (m, n) -closed ideal of R/I if and only if J is an (m, n) -closed ideal of R .*

Proof. Note that the mapping $f : R \rightarrow R/I$ defined by $f(r) = r + I$ is a ring epimorphism. In addition, for this epimorphism f , we also have $f(J) = J/I$ and $f^{-1}(J/I) = J$. Let $u \in \ker(f)$, then $0_R + I = f(x) = x + I$. Consequently, we have $x = x - 0_R \in I \subseteq J$. This means that $\ker(f) \subseteq J$. Finally, by applying Theorem 2.10 on [5], the proof is done. $\square$

The first result is the addition $I_1 + I_2 + \cdots + I_p$ is (m, n) -closed ideal of R if $I_1, I_2, \dots, I_p$ are (m, n) -closed ideals of R .

Theorem 3.5. *Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $I_1, I_2, \dots, I_p$ are (m, n) -closed ideals of R , then $I_1 + I_2 + \cdots + I_p$ is an (m, n) -closed ideal of R .*

Proof. We will prove by induction on p . Let $p = 2$, i.e. I_1 and I_2 are (m, n) -closed ideals of R . Define $f : I_1 \rightarrow \frac{I_1 + I_2}{I_2}$ by $f(i) = i + I_2$ for every $i \in I_1$, then f is well defined. Moreover, f is a ring homomorphism from I_1 to $\frac{I_1 + I_2}{I_2}$. Let $j \in \ker(f)$, then we have $j \in I_2$. Since $j \in I_1$, we have $\ker(f) \subseteq I_1 \cap I_2$. For every $x \in I_1 \cap I_2$, we see that $x = x + I_2 = 0_R + I_2$ and moreover, $x \in \ker(f)$. This means that $I_1 \cap I_2 \subseteq \ker(f)$. Next, we also have $\text{im}(f) = \frac{I_1 + I_2}{I_2}$. It follows from the fundamental theorem of ring homomorphism that $\frac{I_1}{I_1 \cap I_2} \cong \frac{I_1 + I_2}{I_2}$. By Corollary 3.4, $\frac{I_1}{I_1 \cap I_2}$ is an (m, n) -closed ideal in $R/(I_1 \cap I_2)$. Using the isomorphism that we have just proven, $\frac{I_1 + I_2}{I_2}$ is an (m, n) -closed ideal of R/I_2 . Using Corollary 3.4 again, $I_1 + I_2$ is an (m, n) -closed ideal of R . Assume that the theorem is true for $p = k$ and we will prove that it is also true for $p = k + 1$. Let $J = I_1 + I_2 + \cdots + I_k$. By using same method for $p = 2$ and induction hypothesis, ideal $J + I_{k+1} = I_1 + I_2 + \cdots + I_k + I_{k+1}$ is an (m, n) -closed ideal of R . $\square$

Not only Theorem 3.5, Corollary 3.4 is also useful to prove the following theorem.

Theorem 3.6. *Given a proper ideal I of a commutative ring R with identity 1_R and m, n are natural numbers such that $1 \leq n < m$. Ideal $I + \langle x \rangle$ is an (m, n) -closed ideal of $R[x]$ if and only if I is an (m, n) -closed ideal of R .*

Proof. It is sufficient to prove that $\frac{I+\langle x \rangle}{I} \cong I$. Define an $f : I + \langle x \rangle \rightarrow I$ by $f(i + p(x)) = i$ for every $i + p(x) \in I + \langle x \rangle$. Thus, f is well defined and is also a ring homomorphism. Let $y + h(x) \in \ker(f)$, then we have $y = f(y + p(x)) = 0_R$. Moreover, we get $y + p(x) = p(x) \in \langle x \rangle$. This means $\ker(f) \subseteq \langle x \rangle$. Let $f(x) \in \langle x \rangle$, then $r(x) = u(x)x$ for some $u(x) \in R[x]$. Applied it to f , $f(r(x)) = f(0_R + r(x)) = 0_R$. Consequently, $\langle x \rangle$ is a subset of $\ker(f)$. By the definition of f , we have $\text{im}(f) = I$. It follows from the fundamental theorem of ring homomorphism that $\frac{I+\langle x \rangle}{\langle x \rangle} \cong I$. Using Corollary 3.4, $I + \langle x \rangle$ is an (m, n) -closed ideal of $R[x]$ if and only if $\frac{I+\langle x \rangle}{\langle x \rangle}$ is an (m, n) -closed ideal of $\frac{R[x]}{\langle x \rangle}$. Using the fact $\frac{I+\langle x \rangle}{\langle x \rangle} \cong I$, the proof is completely done. $\square$

Let $\{R_i | i = 1, 2, \dots, k\}$ be a finite collection of commutative rings R_i with identity 1_{R_i} . A direct product $\prod_{i=1}^k R_i$ forms commutative ring with identity $(1_{R_1}, 1_{R_2}, \dots, 1_{R_k})$ under componentwise addition and multiplication operations. If I_i is a proper ideal of R_i , for every $i = 1, 2, \dots, k$, then $\prod_{i=1}^k I_i$ is a proper ideal of $\prod_{i=1}^k R_i$. Now, we extend Theorem 2.12 on [5].

Theorem 3.7. *Let R_i be commutative ring with identity 1_{R_i} for every $i = 1, 2, \dots, k$. If I_i is an (m_i, n_i) -closed ideal of R_i for every $i = 1, 2, \dots, k$, then $\prod_{i=1}^k I_i$ is an (m, n) -closed ideal of $\prod_{i=1}^k R_i$ for every natural numbers $m \leq \min\{m_1, m_2, \dots, m_k\}$ and $n \geq \max\{n_1, n_2, \dots, n_k\}$.*

Proof. Let $y \in \prod_{i=1}^k R_i$ such that $y^m \in \prod_{i=1}^k I_i$. Then y^m can be expressed as $y^m = (r_1, r_2, \dots, r_k)^m = (r_1^m, r_2^m, \dots, r_k^m) \in \prod_{i=1}^k I_i$. This means that $r_i^m \in I_i$ for every i . Furthermore, we have $r_i^{m_i} = r_i^{m_i-m} r_i^m \in I_i$ for every i . Note that for any $i = 1, 2, \dots, k$, I_i is an (m_i, n_i) -closed ideal, so we get $r_i^{n_i} \in I_i$ for every $i = 1, 2, \dots, k$. Moreover, it is obvious that $r_i^n = r_i^{n-n_i} r_i^{n_i} \in I_i$ for every $i = 1, 2, \dots, k$. Finally, we get $y^n = (r_1^n, r_2^n, \dots, r_k^n) = (r_1, r_2, \dots, r_k)^n \in \prod_{i=1}^k I_i$. $\square$

The zero ideal $\{0_R\}$ is a prime ideal of an integral domain R , since integral domain does not have zero divisor elements. Consequently, ideal $\{0_R\}$ is a semiprime ideal of an integral domain R . Moreover, we have $\{0_R\}$ is an (m, n) -closed ideal of an integral domain R for every natural numbers m and n such that $1 \leq n < m$. Finally, we have an immediate consequence of Theorem 2.10 on [5].

Corollary 3.8. *Let R be a commutative ring with identity 1_R and S be an integral domain. If $f : R \rightarrow S$ is a ring homomorphism, then $\ker(f)$ is an (m, n) -closed ideal of R for every natural numbers m and n such that $1 \leq n < m$.*

Proof. Using Theorem 2.10 on [5] and $\{0_S\}$ is an (m, n) -closed ideal for every $1 \leq n < m$, $f^{-1}(\{0_S\}) = \ker(f)$ is an (m, n) -closed ideal for every natural numbers m and n such that $1 \leq n < m$. $\square$

We give an example to understand Corollary 3.8.

Example 3.9. Consider a ring homomorphism $f : \mathbb{Z} \times \mathbb{Z} \rightarrow \mathbb{Z}$ defined by $f((a, b)) = a$ for every $(a, b) \in \mathbb{Z} \times \mathbb{Z}$. Using Corollary 3.8, we get

$$\begin{aligned} \ker(f) &= \{(a, b) \in \mathbb{Z} \times \mathbb{Z} \mid a = f((a, b)) = 0\} \\ &= \{(0, b) \mid b \in \mathbb{Z}\} = \langle (0, 1) \rangle \end{aligned}$$

is an (m, n) -closed ideal of $\mathbb{Z} \times \mathbb{Z}$ for every natural numbers m and n such that $1 \leq n < m$.

3.2. Quasi (m, n) -closed Ideal. Let m and n be natural numbers such that $1 \leq n < m$ and R be a commutative ring with identity 1_R . Define three collections,

$$\begin{aligned} \mathcal{A}_{(m,n)} &= \{I \mid I \text{ is an } (m, n) \text{-closed ideal of } R\}, \\ \mathcal{A}'_{(m,n)} &= \{I \mid I \text{ is an } (m, m-n) \text{-closed ideal of } R\}, \text{ and} \\ \mathcal{B}_{(m,n)} &= \{I \mid I \text{ is a quasi } (m, n) \text{-closed ideal of } R\}. \end{aligned}$$

Note that $\mathcal{A}_{(m,n)} \neq \emptyset$ since $\text{nil}(R)$ is an (m, n) -closed ideal. This is also true for $\mathcal{A}'_{(m,n)}$. Since a quasi (m, n) -closed ideal is a generalization from (m, n) -closed ideal and $(m, m-n)$ -closed ideal, then $\mathcal{A}_{(m,n)} \subseteq \mathcal{B}_{(m,n)}$ and $\mathcal{A}'_{(m,n)} \subseteq \mathcal{B}_{(m,n)}$. It is easy to prove that $\mathcal{A}_{(m,n)} = \mathcal{B}_{(m,n)}$ if $m \leq 2n$ and $\mathcal{A}'_{(m,n)} = \mathcal{B}_{(m,n)}$ if $m \geq 2n$.

From the definition of quasi (m, n) -closed ideal, a quasi (m, n) -closed ideal is a quasi $(m, m-n)$ -closed ideal and vice versa. However, this fact need not to be true for (m, n) -closed ideal. Example 2.2 on [5] shows that $\langle 16 \rangle$ is semi 2-absorbing (i.e. $(3, 2)$ -closed) ideal of $\mathbb{Z}$. Unfortunately, if we choose $x = 4 \in \mathbb{Z}$, then we have $4^3 \in \langle 16 \rangle$ and $4 \notin \langle 16 \rangle$. Hence, $\langle 16 \rangle$ is not a $(3, 1)$ -closed ideal of $\mathbb{Z}$. Conversely, an $(m, m-n)$ -closed ideal need not to be true an (m, n) -closed ideal. Ideal $\langle 8 \rangle$ is a $(4, 4-1)$ -closed ideal of $\mathbb{Z}$, since for every $x \in \mathbb{Z}$ with $x^4 \in \langle 8 \rangle$, then $x \in \sqrt[4]{\langle 8 \rangle} = \langle 2 \rangle$ and so that $x^3 \in \langle 8 \rangle$. However, there is $2 \in \mathbb{Z}$ with $2^4 \in \langle 8 \rangle$ and $2 \notin \langle 8 \rangle$. Hence, $\langle 8 \rangle$ is not a $(4, 1)$ -closed ideal of $\mathbb{Z}$.

On the above paragraph, we have $\mathcal{A}_{(m,n)} = \mathcal{B}_{(m,n)}$ if $m \leq 2n$ and $\mathcal{A}'_{(m,n)} = \mathcal{B}_{(m,n)}$ if $m \geq 2n$. Hence, those facts make an immediate consequence.

Corollary 3.10. Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $m = 2n$, then the following statements are equivalent.

1. A proper ideal I of R is an (m, n) -closed ideal of R .
2. A proper ideal I of R is an $(m, m-n)$ -closed ideal of R .
3. A proper ideal I of R is a quasi (m, n) -closed ideal of R .

Recall that proper ideals I_1 and I_2 of a commutative ring R with identity 1_R are comaximal if $I_1 + I_2 = R$. If the collection of ideals of R $\{I_i \mid i = 1, 2, \dots, k\}$ are pairwise comaximal, then $I_1 \cap I_2 \cap \dots \cap I_k = I_1 I_2 \dots I_k$ [14]. Now, we modify Corollary 2.4 on [5] in terms of quasi (m, n) -closed ideal.

Theorem 3.11. Let R be a commutative ring with identity 1_R , m and n be natural numbers such that $1 \leq n < m$ and $I_1, I_2, \dots, I_k$ be quasi (m, n) -closed ideals, then we have :

1. The intersection $I_1 \cap I_2 \cap \cdots \cap I_k$ is also a quasi (m, n) -closed ideal of R .
2. If $I_1, I_2, \dots, I_k$ are pairwise comaximal, then $I_1 I_2 \cdots I_k$ is a quasi (m, n) -closed ideal of R .

Proof. Since Corollary 2.4 on [5] holds for both cases $m \leq 2n$ and $m \geq 2n$, then the first statement is fulfilled. The second statement immediately follows from the first statement. $\square$

Theorem 3.1 can be used to prove the following theorem.

Theorem 3.12. *Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $\{I_\alpha | \alpha \in \Lambda\}$, where $\emptyset \neq \Lambda$ denotes an indexing set, is a collection of quasi (m, n) -closed ideals of R , then $\cap_{\alpha \in \Lambda} I_\alpha$ is a quasi (m, n) -closed ideal of R .*

Proof. Since Theorem 3.1 holds for both cases $m \leq 2n$ and $m \geq 2n$, then it is true that $\cap_{\alpha \in \Lambda} I_\alpha$ is a quasi (m, n) -closed ideal of R . $\square$

Let a be an element of a commutative ring R with identity 1_R . If I is an ideal of R , then $I_a = \{x \in R | ax \in I\}$ is an ideal of R . Recall that $a \in R$ is an idempotent element if $a^2 = a$ [12].

Theorem 3.13. *Let I be a proper ideal of a commutative ring R with identity 1_R , m and n be natural numbers such that $1 \leq n < m$ and $a \in R - I$ be a nonunit idempotent element of R . If I is a quasi (m, n) -closed ideal of R , then I_a is a quasi (m, n) -closed ideal of R .*

Proof. The condition $a \in R - I$ be a nonunit element of R implies that I_a is a proper ideal of R . Let $x \in R$ such that $x^m \in I_a$. Assume that $x^{m-n} \notin I_a$. Consequently, we have $(ax)^m = ax^m \in I$ and $(ax)^{m-n} = ax^{m-n} \notin I$. Since I is a quasi (m, n) -closed ideal of R , then $(ax)^n = ax^n \in I$. This means that $x^n \in I_a$. $\square$

The following theorem provide a necessary condition for a proper ideal I of R be a quasi (m, n) -closed ideal of R .

Theorem 3.14. *Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If whenever $J^m \subseteq I$ for every ideal J of R implies $J^n \subseteq I$ or $J^{m-n} \subseteq I$, then I is a quasi (m, n) -closed ideal of R .*

Proof. Let $x \in R$ such that $x^m \in I$. If $\langle x \rangle$ is an ideal generated by x , then $\langle x \rangle^m \subseteq I$. By assumption, we have $\langle x \rangle^n \subseteq I$ or $\langle x \rangle^{m-n} \subseteq I$. If $\langle x \rangle^n \subseteq I$, then $x^n \in \langle x \rangle^n \subseteq I$. If $\langle x \rangle^{m-n} \subseteq I$, then $x^{m-n} \in \langle x \rangle^{m-n} \subseteq I$. $\square$

Corollary 2.11 on [5] can be modified for quasi (m, n) -closed ideal. It appears in Corollary 2 on [6]. Hence, Corollary 3.3 can be extended in terms of quasi (m, n) -closed. We omit its proof since it is very similar to Corollary 3.3.

Corollary 3.15. *Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $I[x]$ is a quasi (m, n) -closed ideal of $R[x]$, then I is a quasi (m, n) -closed ideal of R .*

Corollary 3.4 also holds for quasi (m, n) -closed ideal as we can see in Corollary 2 on [6]. Hence, we can extended Theorem 3.5 and Theorem 3.6 in terms of quasi (m, n) -closed ideal. Same with above corollary, we omit their proofs.

Theorem 3.16. *Let R be a commutative ring with identity 1_R and m, n be natural numbers such that $1 \leq n < m$. If $I_1, I_2, \dots, I_p$ are quasi (m, n) -closed ideals of R , then $I_1 + I_2 + \dots + I_p$ is a quasi (m, n) -closed ideal of R .*

Theorem 3.17. *Given a proper ideal I of a commutative ring R with identity 1_R and m, n are natural numbers such that $1 \leq n < m$. Ideal $I + \langle x \rangle$ is a quasi (m, n) -closed ideal of $R[x]$ if and only if I is a quasi (m, n) -closed ideal of R .*

Now, we present Corollary 3 on [6] without proof to support our next result.

Corollary 3.18. *Let I_i be an ideal of commutative ring R_i with identity 1_{R_i} for every $i = 1, 2, \dots, k$. Let m and n are natural numbers such that $1 \leq n < m$. Then we have :*

1. *If I_i is a proper ideal of R_i for some $i = 1, 2, \dots, k$, then I_i is a quasi (m, n) -closed ideal of R_i if and only if $\prod_{j=1}^{i-1} R_i \times I_i \times \prod_{j=i+1}^k R_i$ is a quasi (m, n) -closed ideal of $\prod_{i=1}^k R_i$.*
2. *If I_i is a quasi (m_i, n_i) -closed ideal of R_i for any $i = 1, 2, \dots, k$ and $t = \max\{n_i, m_i - n_i | i = 1, 2, \dots, k\}$, then $\prod_{i=1}^k I_i$ is a quasi (m, n) -closed ideal of $\prod_{i=1}^k R_i$ whenever $m \leq \min\{m_1, m_2, \dots, m_k\}$ and $m \geq 2t$.*

Corollary 3.18 can be used to prove our result.

Corollary 3.19. *Let I_i be an ideal of commutative ring R_i with identity 1_{R_i} for every $i = 1, 2, \dots, k$ and n be a natural numbers. Then we have :*

1. *If I_i is a proper ideal of R_i for some $i = 1, 2, \dots, k$, then I_i is a quasi $(2n, n)$ -closed ideal of R_i if and only if $\prod_{j=1}^{i-1} R_i \times I_i \times \prod_{j=i+1}^k R_i$ is a quasi $(2n, n)$ -closed ideal of $\prod_{i=1}^k R_i$.*
2. *If I_i is a quasi $(2n, n)$ -closed ideal of R_i for any $i = 1, 2, \dots, k$, then $\prod_{i=1}^k I_i$ is a quasi $(2n, n)$ -closed ideal of $\prod_{i=1}^k R_i$.*

Proof. Let $m = 2n$ for every $i = 1, 2, \dots, k$. Using Corollary 3.18 and facts that $m = 2n = \min\{2n\}$ and $m = 2n = 2 \max\{n, 2n - n\} = 2 \max\{n\}$, the proof is completely done. $\square$

Example 3.20. *Consider proper ideals of $\mathbb{Z}$, i.e. $\langle 4 \rangle, \langle 9 \rangle$ and $\langle 25 \rangle$. We can check that $\sqrt{\langle 4 \rangle} = \langle 2 \rangle, \sqrt{\langle 9 \rangle} = \langle 3 \rangle$ and $\sqrt{\langle 25 \rangle} = \langle 5 \rangle$. Let $x \in \mathbb{Z}$ such that $x^4 \in \langle 4 \rangle$. Then $x \in \langle 2 \rangle$ and so that $x^2 \in \langle 2 \rangle$. Certainly, this fact also satisfy for $\langle 9 \rangle$ and $\langle 25 \rangle$. Hence, $\langle 4 \rangle, \langle 9 \rangle$ and $\langle 25 \rangle$ are quasi $(2n, n)$ -closed ideals of $\mathbb{Z}$. Using Corollary 3.19, product $\langle 4 \rangle \times \langle 9 \rangle \times \langle 25 \rangle$ is a quasi $(2n, n)$ -closed ideal of $\mathbb{Z} \times \mathbb{Z} \times \mathbb{Z}$.*

We already know that ideal $\{0_R\}$ is an (m, n) -closed of an integral domain R for every natural numbers m and n such that $1 \leq n < m$. Thus, it is a quasi (m, n) -closed ideal of a integral domain R for every natural numbers m and n such that $1 \leq n < m$. Hence, Corollary 3.8 also holds for quasi (m, n) -closed ideal. It can be proved with Proposition 3 on [6].

Corollary 3.21. *Let R be a commutative ring with identity 1_R and S be an integral domain. If $f : R \rightarrow S$ is a ring homomorphism, then $\ker(f)$ is a quasi (m, n) -closed ideal for every natural numbers m and n such that $1 \leq n < m$.*

Recall that a commutative ring R with identity 1_R is called a local ring if R has exactly one maximal ideal. A local ring R with its maximal ideal M is denoted by (R, M) [13]. Hence, we can modify Lemma 2.12 on [7].

Theorem 3.22. *Let m and n be natural numbers such that $1 \leq n < m$. If (R, M) is a local ring which satisfies either $M^n = \{0_R\}$ or $M^{m-n} = \{0_R\}$, then every proper ideal I of R is a quasi (m, n) -closed ideal of R .*

Proof. If $M^n = \{0_R\}$ holds, then by Lemma 2.12 on [7], every proper ideal I is an (m, n) -closed ideal of R . If $M^{m-n} = \{0_R\}$ holds, then by Lemma 2.12 on [7], every proper ideal I is an $(m, m-n)$ -closed ideal of R . Hence, every proper ideal I of R is a quasi (m, n) -closed ideal of R . $\square$

4. CONCLUSION

In this paper, we have presented several new properties about (m, n) -closed ideal and quasi (m, n) -closed ideal. The new properties have been derived from several theorems and corollaries on [5] and [6].

Acknowledgement. We thank the reviewers for their suggestions and revisions to this paper.

REFERENCES

- [1] D. Malik, J. Mordeson, and M. Sen, *Fundamentals of Abstract Algebra*. Mc.Graw-Hill Companies Inc., Unites States of America, 1999. [\[1\]](#)
- [2] D. Anderson and A. Badawi, "On n -absorbing ideals of commutative rings," *Commun. Algebra*, vol. 39, pp. 1646–1672, 2011. [\[2\]](#)
- [3] H. Choi, "On n -absorbing ideals of locally divided commutative rings," *Commun. Algebra*, vol. 54, pp. 483–513, 2022. [\[3\]](#)
- [4] H. Moghimi and S. Naghani, "On n -absorbing ideals and the n -krull dimension of a commutative ring," *J. Korean Math. Soc.*, vol. 53, pp. 1225–1236, 2016. [\[4\]](#)
- [5] D. Anderson and A. Badawi, "On (m, n) -closed ideals of commutative rings," *J. Algebra Appl.*, vol. 16, pp. 1–21, 2017. [\[5\]](#)

- [6] H. Khashan and E. Celikel, “A new generalization of (m, n) -closed ideals,” *J. Math. Sci.*, vol. 280, pp. 288–299, 2024. [\[1\]](#)
- [7] M. Issoual, M. Mahdou, and M. Moutui, “On (m, n) -closed ideals in amalgamated algebra,” *Int. Electron. J. Algebra*, vol. 29, pp. 134–147, 2021. [\[1\]](#)
- [8] A. Badawi, M. Issoual, and M. Mahdou, “On n -absorbing ideals and (m, n) -closed ideals in trivial ring extension of commutative rings,” *J. Algebra Appl.*, vol. 18, pp. 1–19, 2019. [\[1\]](#)
- [9] D. Anderson and M. Bataineh, “Generalizations of prime ideals,” *Commun. Algebra*, vol. 36, pp. 686–696, 2008. [\[1\]](#)
- [10] A. Badawi, “On 2-absorbing ideals of commutative rings,” *Bull. Aust. Math. Soc.*, vol. 75, pp. 417–429, 2007. [\[1\]](#)
- [11] H. Mostafanasab and A. Darani, “On n -absorbing ideals and two generalizations of semiprime ideals,” *Thai J. Math.*, vol. 15, pp. 272–794, 2017. [\[1\]](#)
- [12] D. Dummit and R. Foote, *Abstract Algebra Third Edition*. John Wiley and Sons Inc., 2004. [\[1\]](#)
- [13] B. Singh, *Basic Commutative Algebra*. World Scientific Publishing Co.Pte. Ltd., 2011. [\[1\]](#)
- [14] R. Sharp, *Steps in Commutative Algebra Second Edition*. Cambridge University Press, 2000. [\[1\]](#)

On Graded Strongly 1-Absorbing Primary Ideals and Graded 2-Absorbing I-Primary Ideals

Maria Asa Pitayani^{1*}, Sutopo², Sri Wahyuni³

^{1,2,3}Department of Mathematics, Universitas Gadjah Mada, Indonesia

¹mariaasapitayani@mail.ugm.ac.id, ²sutopo.mipa@ugm.ac.id, ³swahyuni@ugm.ac.id

Abstract. Given a group G with identity e and a G -graded commutative ring R with unity element 1_R . This paper introduces a new concept, namely, graded strongly 1-absorbing primary ideals, which is a subclass of the graded 1-absorbing primary ideals. A proper graded ideal I of R is said to be a graded strongly 1-absorbing primary ideal of R if whenever non-unit elements $a, b, c \in h(R)$ with $abc \in I$, then either $ab \in I$ or $c \in \text{Grad}(0)$. Several properties of graded strongly 1-absorbing primary ideals will be investigated in this paper. Furthermore, a new structure called graded 2-absorbing I -primary ideal is also introduced.

Keywords: Graded 1-Absorbing Primary Ideal, Graded Strongly 1-Absorbing Primary Ideal, Graded 2-Absorbing I -Primary Ideal.

1. INTRODUCTION

Throughout this article, G will be a group with identity e and R will be an abelian ring with a nonzero unity 1_R . A ring R is said to be G -graded if there exists a family of additive subgroups $\{R_g\}_{g \in G}$ such that $R = \bigoplus_{g \in G} R_g$ and $R_{g_1}R_{g_2} \subseteq R_{g_1+g_2}$ for every $g_1, g_2 \in G$. If $R_{g_1}R_{g_2} = R_{g_1+g_2}$ for all $g_1, g_2 \in G$, then R is said to be a strongly graded ring. Moreover, a G -graded ring R is positively graded if $R_g = \{0\}$ for all $g < 0$ and negatively graded if $R_g = \{0\}$ for all $g > 0$. The set $h(R) = \bigcup_{g \in G} R_g$ is the set of homogeneous elements of R . A nonzero element $a_g \in R_g$ is said to be a homogeneous element of degree g , and can be written as $\deg(a_g) = g$. Every nonzero element in a G -graded ring R can be uniquely expressed as a finite sum of homogeneous elements, denoted by $a = \sum_{g \in G} a_g$, where a_g is a homogeneous component of a in R_g .

*Corresponding author

2020 Mathematics Subject Classification: 11R44

Received: 24-05-2025, accepted: 02-08-2025.

Let S be a subring of a G -graded ring R . Then S is said to be a graded subring if $S = \bigoplus_{g \in G} S_g$, where $S_g = S \cap R_g$ for each $g \in G$. Since S is a graded subring of R , then $S_{g_1} S_{g_2} = (S \cap R_{g_1})(S \cap R_{g_2}) \subseteq S \cap R_{g_1+g_2} = S_{g_1+g_2}$ for all $g_1, g_2 \in G$. Consequently, every graded subring S of R is a G -graded ring.

Analogous to the concept of an ideal in a ring, the concept of a graded ideal is introduced in graded rings. An ideal I of a graded ring R is said to be a graded ideal if I respects the grading structure, i.e., $I = \bigoplus_{g \in G} I_g$, where $I_g = I \cap R_g$ for each $g \in G$. An important construction related to graded ideals is the graded radical of a graded ideal I , denoted by $\text{Grad}(I)$. The graded radical consists of all elements $a = \sum_{g \in G} a_g \in R$ such that for every $g \in G$, there exists a positive integer n_g satisfying $a_g^{n_g} \in I$, denoted by

$$\text{Grad}(I) = \left\{ a = \sum_{g \in G} a_g \in R \mid \forall g \in G, \exists n_g \in \mathbb{N}, \ni a_g^{n_g} \in I \right\}.$$

It has been established that the graded radical of a graded ideal is a graded ideal of R .

The following summarizes several important properties related to graded ideals in graded rings.

Lemma 1.1. [1] *Let R be a G -graded ring. The following hold:*

- (1) *If I and J are graded ideals in R , then $I + J$, IJ , and $I \cap J$ are graded ideals in R .*
- (2) *If $a \in h(R)$, then Ra is a graded ideal in R .*

Lemma 1.2. [2] *Let R be a G -graded ring, and let I, J be graded ideals in R . Then,*

- (1) $\text{Grad}(\text{Grad}(I)) = \text{Grad}(I)$
- (2) $\text{Grad}(IJ) = \text{Grad}(I \cap J) = \text{Grad}(I) \cap \text{Grad}(J)$.

Proof. It follows from ([2], Proposition 2.4). □

Refai, in [3], introduced a generalization of the graded prime ideal, called the graded primary ideal. A graded ideal I of a G -graded ring R is said to be a graded primary ideal if $I \neq R$ and whenever $a, b \in h(R)$ with $ab \in I$, then $a \in I$ or $a \in \text{Grad}(I)$. Every graded prime ideal is a graded primary ideal, but the converse does not generally hold. Later, in [4], a further generalization of the graded primary ideal, called the graded 2-absorbing primary ideal, was studied. A proper graded ideal I of R is said to be a graded 2-absorbing primary ideal if whenever $a, b, c \in h(R)$ with $abc \in I$, then $ab \in I$ or $ac \in \text{Grad}(I)$ or $bc \in \text{Grad}(I)$. In [5], Abu-Dawwas and Bataineh introduced a new subclass of graded 2-absorbing primary ideals, called graded 1-absorbing primary ideals. A proper graded ideal I of a graded ring R is said to be a graded 1-absorbing primary if whenever non-unit elements $a, b, c \in h(R)$ such that $abc \in I$, then $ab \in I$ or $c \in \text{Grad}(I)$. Since the concept of a graded 1-absorbing primary ideal generalizes that of a graded primary ideal, it follows that every graded primary ideal is necessarily a graded

1-absorbing primary ideal. However, the converse is not necessarily true; that is, a graded 1-absorbing primary ideal is not always a graded primary ideal.

In 2015, [6] studied rings in which every non-unit element is a product of a unit and a nilpotent element, referring to them as UN rings. Building upon this concept, the structure of the UN rings was later extended to graded rings.

Definition 1.3. [7] A G -graded ring R is called a HUN ring if every homogeneous element in R is either a unit or nilpotent.

Example 1.4. [8] Let R be a graded field, and let $u \notin R$ be an element such that $u^2 = 1$. Define a graded field $F = \{a + ub \mid a, b \in R \text{ and } u^2 = 1\}$ with respect to the group $\mathbb{Z}_2$, where the grading is given by $F_0 = R$ and $F_1 = uR$. Next, we prove that F is a HUN-ring. Let $a \in F_0$. Since every element of the field R is a unit, it follows that a is a unit in F_0 . Now, consider the elements in F_1 . Take any $ub \in F_1$ with $u^2 = 1$. Since $b \in R$ and b is a unit, we obtain $(ub)^2 = u^2b^2 = b^2$. Consequently, b^2 is trivially nilpotent if $b = 0$. Conversely, if $b \neq 0$, then b^2 is a unit. Therefore, the graded field F is a HUN ring.

Previously, in [9], Almahdi et al. introduced the concept of a strongly 1-absorbing primary ideal, which can be used to characterize UN rings and local rings, rings that have exactly one maximal ideal. Based on this idea, in [8], Abu-Dawwas developed a similar concept in graded rings, introducing the graded strongly 1-absorbing primary ideal, which forms a new subclass of graded 1-absorbing primary ideals. Just like the strongly 1-absorbing primary ideal, the graded strongly 1-absorbing primary ideal can also be used to characterize HUN rings and graded local rings.

This paper is a review of the work previously written by Abu-Dawwas in [8]. Our contributions in this paper include providing examples, adding properties, and deriving corollaries from the existing propositions. We give examples (see Example 2.2) of a graded strongly 1-absorbing primary ideal, and derive corollaries from Proposition 2.12, particularly when R is a graded Noetherian. Furthermore, we investigate the structure of graded ideals in the direct product of graded rings. Since the graded ideal in this ring is not a graded strongly 1-absorbing primary ideal, we examine whether it satisfies the properties of a graded 2-absorbing I -ideal, which is a generalization of graded prime ideals and was introduced by I. Akray, Adil K. Jabbar and Shadan A. Othman in [10].

2. GRADED STRONGLY 1-ABSORBING PRIMARY IDEAL

In this section, we study the concept of graded strongly 1-absorbing primary ideals, a subclass of graded 1-absorbing primary ideals.

Definition 2.1. [8] Let R be a G -graded ring and I a proper graded ideal of R . The ideal I is said to be a graded strongly 1-absorbing primary ideal in R if whenever non-unit elements $a, b, c \in h(R)$ such that $abc \in I$, then either $ab \in I$ or $c \in \text{Grad}(\{0\})$.

The following example illustrates the concept of a graded strongly 1-absorbing primary ideal.

Example 2.2. Let $R = \mathbb{R}[x]/\langle x^9 \rangle$ be a $\mathbb{Z}$ -graded ring and consider the graded ideal $I = \langle x^3 \rangle$ in R . Then I is a graded strongly 1-absorbing primary ideal. Let $\overline{p(x)} = \sum_{i=0}^8 a_i x^i$, $\overline{q(x)} = \sum_{i=0}^8 b_i x^i$, $\overline{r(x)} = \sum_{i=0}^8 c_i x^i \in h(R)$ be non-unit elements such that $\overline{p(x)q(x)r(x)} \in I$. If $\overline{p(x)q(x)} \in I$, then I is trivially a graded strongly 1-absorbing primary ideal. Suppose $\overline{p(x)q(x)} \notin I$. Then we analyze the following cases:

- (1) If $\deg(\overline{p(x)q(x)}) = 0$, then $\deg(\overline{r(x)}) \geq 3$, so there exists $n_i \geq 3$ such that $(\overline{r(x)})^{n_i} \in \{\overline{0}\}$. Hence, $\overline{r(x)} \in \text{Grad}(\{\overline{0}\})$.
- (2) If $\deg(\overline{p(x)q(x)}) = 1$, then $\deg(\overline{r(x)}) \geq 2$, so there exists $n_i \geq 5$ such that $(\overline{r(x)})^{n_i} \in \{\overline{0}\}$. Hence, $\overline{r(x)} \in \text{Grad}(\{\overline{0}\})$.
- (3) If $\deg(\overline{p(x)q(x)}) = 2$, then $\deg(\overline{r(x)}) \geq 1$, so there exists $n_i \geq 9$ such that $(\overline{r(x)})^{n_i} \in \{\overline{0}\}$. Hence, $\overline{r(x)} \in \text{Grad}(\{\overline{0}\})$.

Thus, for any $\overline{p(x)}, \overline{q(x)}, \overline{r(x)} \in h(R)$ such that $\overline{p(x)q(x)r(x)} \in I$, we have either $\overline{p(x)q(x)} \in I$ or $\overline{r(x)} \in \text{Grad}(\{\overline{0}\})$. Therefore, I is a graded strongly 1-absorbing primary ideal in R . In general, the ideal $I = \langle x^p \rangle$ is a graded strongly 1-absorbing primary ideal of $R = \mathbb{R}[x]/\langle x^p \rangle$.

As mentioned earlier, the concept of a graded strongly 1-absorbing primary ideal forms a new subclass of graded 1-absorbing primary ideals. Hence, every graded strongly 1-absorbing primary ideal is a graded 1-absorbing primary ideal, but the converse does not necessarily hold. Not every graded 1-absorbing primary ideal is graded strongly 1-absorbing. The following example illustrates this.

Example 2.3. Let $R = \mathbb{Z}[i]$ and $G = \mathbb{Z}_2$. Then R be a G -graded ring by $R_0 = \mathbb{Z}$ and $R_1 = i\mathbb{Z}$. Consider the graded ideal $I = 3R$, which is a graded 1-absorbing primary ideal. For arbitrary non-unit elements $a, b \in h(R)$ with $ab \in I$ and $a \notin I$, it follows that 3 divides ab but does not divide a . Thus, 3 must divide b , which means $b \in \text{Grad}(I)$. Hence, I is a graded 1-absorbing primary ideal. However, I is not a graded strongly 1-absorbing primary ideal. Specifically, consider $2, 3 \in h(R)$ such that $2 \cdot 2 \cdot 3 \in I$, but $2 \cdot 2 \notin I$ and $3 \notin \text{Grad}(\{0\})$. Therefore, we conclude that I is a graded 1-absorbing primary ideal but not a graded strongly 1-absorbing primary ideal.

Not every G -graded ring necessarily contains a graded strongly 1-absorbing primary ideal. However, a G -graded ring does contain a graded strongly 1-absorbing primary ideal if it has the graded prime ideal $\text{Grad}(\{0\})$ or is a graded local ring. This result is formalized in the following theorem.

Theorem 2.4. [8] Let R be a G -graded ring. A graded strongly 1-absorbing primary ideal exists in R if and only if

- (1) the ideal $\text{Grad}(0)$ is a graded prime ideal, or
- (2) the ring R is a graded local ring.

Proof. It follows from ([8], Theorem 2.8.). $\square$

In ring theory, if R and S are rings, then the direct product $R \times S$ forms a ring. Similarly, if R and S are G -graded rings with grading $\{R_g\}_{g \in G}$ and $\{S_g\}_{g \in G}$, then the direct product $R \times S$ is also a G -graded ring with grading $(R \times S)_g = R_g \times S_g$ for each $g \in G$.

Corollary 2.5. [8] *If R and S are G -graded rings, then $R \times S$ does not contain any graded strongly 1-absorbing primary ideal.*

Proof. Let R and S be G -graded rings. If neither R nor S is a graded local ring, then it is clear that $R \times S$ is also not a graded local ring. Suppose that R and S are graded local rings with graded maximal ideals $\mathcal{M}_R$ and $\mathcal{M}_S$, respectively. Consequently, $R \times S$ has more than one graded maximal ideal, namely $\mathcal{M}_R \times S$ and $R \times \mathcal{M}_S$. Therefore, $R \times S$ is not a graded local ring. Next, we examine whether $\text{Grad}(\{0_{R \times S}\}) = \text{Grad}(\{0_R\}) \times \text{Grad}(\{0_S\})$ is a graded prime ideal in $R \times S$. Consider elements $(a, b) \notin \text{Grad}(\{0_{R \times S}\})$ where $a \in \text{Grad}(\{0_R\})$ and $b \notin \text{Grad}(\{0_S\})$, as well as $(c, d) \notin \text{Grad}(\{0_{R \times S}\})$ where $c \notin \text{Grad}(\{0_R\})$ and $d \in \text{Grad}(\{0_S\})$. However, since there exist $m, n \in \mathbb{N}$ such that $(ac)^n = a^n c^n = 0_R c^n = 0_R$ and $(bd)^m = b^m d^m = b^m 0_S = 0_S$, it follows that $(ac, bd) \in \text{Grad}(\{0_{R \times S}\})$. As a result, $\text{Grad}(\{0_{R \times S}\})$ is not a graded prime ideal in $R \times S$. Since $R \times S$ is not a graded local ring and $\text{Grad}(\{0_{R \times S}\})$ is not a graded prime ideal, we conclude that $R \times S$ does not contain any graded strongly 1-absorbing primary ideal. $\square$

Some properties of graded strongly 1-absorbing primary ideals are presented in the following propositions.

Proposition 2.6. [8] *Let R be a G -graded ring, and let I and J be proper graded ideals of R . If I and J are graded strongly 1-absorbing primary ideals, then $I \cap J$ is also a graded strongly 1-absorbing primary ideal.*

Proof. It follows from ([8], Proposition 2.12). $\square$

The following theorem is a development of the results studied in [11], [3], which were previously investigated in the context of 1-absorbing primary ideals and graded primary ideals. In this study, we further develop these results within the framework of graded strongly 1-absorbing primary ideals.

Definition 2.7. *Let R be a G -graded ring and let I be a graded strongly 1-absorbing primary ideal of R . Then $J = \text{Grad}(I)$ is a graded prime ideal of R , and we say that I is a J -graded strongly 1-absorbing primary ideal.*

So we have the following result.

Proposition 2.8. *Let R be a G -graded ring, and let $I_1, I_2, \dots, I_n$ be proper graded ideal of R . If $I_1, I_2, \dots, I_n$ are J -graded strongly 1-absorbing primary ideal, then $I = \bigcap_{i=1}^n I_i$ is a J -graded strongly 1-absorbing primary ideal.*

Proof. First, we will prove that $\text{Grad}(I) = J$. Let $I_1, I_2, \dots, I_n$ are J -graded strongly 1-absorbing primary ideal. It is given that $\text{Grad}(I_i) = J$ for every $i = 1, 2, \dots, n$. By Proposition 1.2 we have $\text{Grad}(I) = \text{Grad}(\bigcap_{i=1}^n I_i) = \bigcap_{i=1}^n \text{Grad}(I_i)$. Since $\text{Grad}(I_i) = J$ for all i , it follows that $\bigcap_{i=1}^n \text{Grad}(I_i) = J \cap J \cap \dots \cap J = J$. Thus, we have shown that $\text{Grad}(I) = J$. Next, by Proposition 2.6, we know that $I = \bigcap_{i=1}^n I_i$ is a graded strongly 1-absorbing primary ideal of R . Therefore, we conclude that I is a J -graded strongly 1-absorbing primary ideal. $\square$

Proposition 2.9. [8] *Let R be a G -graded ring. If every element of $h(R)$ is either nilpotent or a unit, then Rw is a graded strongly 1-absorbing primary ideal in R for every non-unit element $w \in h(R)$.*

Proof. It follows from ([8], Proposition 2.13). $\square$

Corollary 2.10. [8] *Let R be a graded ring. If R is an HUN ring, then every proper graded ideal in R is a graded strongly 1-absorbing primary ideal.*

Proof. Let R be a HUN ring and I a proper graded ideal in R . By Definition 1.3, every element in $h(R)$ is either nilpotent or a unit. Suppose there exist nonunit elements $a, b, c \in h(R)$ such that $abc \in I$ and $c \notin \text{Grad}(\{0\})$. The objective is to show that every proper graded ideal I in R is a graded strongly 1-absorbing primary ideal. Since $a, b, c \in h(R)$ are nonunit elements and $abc \in R(abc)$, Proposition 2.9 ensures that $R(abc)$ is a graded strongly 1-absorbing primary ideal in R . Consequently, $ab \in R(abc) \subseteq I$. Therefore, it is established that if every element in $h(R)$ is either nilpotent or a unit, then every proper graded ideal in R is a graded strongly 1-absorbing primary ideal. $\square$

Example 2.11. *Consider the ring $\mathbb{Z}/9\mathbb{Z}$. The following table demonstrates that every element in $\mathbb{Z}/9\mathbb{Z}$ is either a unit or nilpotent.*

TABLE 1. Multiplication $(\cdot)$ in $\mathbb{Z}/9\mathbb{Z}$.

$\cdot$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$	$\bar{7}$	$\bar{8}$
$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$	$\bar{0}$
$\bar{1}$	$\bar{0}$	$\bar{1}$	$\bar{2}$	$\bar{3}$	$\bar{4}$	$\bar{5}$	$\bar{6}$	$\bar{7}$	$\bar{8}$
$\bar{2}$	$\bar{0}$	$\bar{2}$	$\bar{4}$	$\bar{6}$	$\bar{8}$	$\bar{1}$	$\bar{3}$	$\bar{5}$	$\bar{7}$
$\bar{3}$	$\bar{0}$	$\bar{3}$	$\bar{6}$	$\bar{0}$	$\bar{3}$	$\bar{6}$	$\bar{0}$	$\bar{3}$	$\bar{6}$
$\bar{4}$	$\bar{0}$	$\bar{4}$	$\bar{8}$	$\bar{3}$	$\bar{7}$	$\bar{2}$	$\bar{6}$	$\bar{1}$	$\bar{5}$
$\bar{5}$	$\bar{0}$	$\bar{5}$	$\bar{1}$	$\bar{6}$	$\bar{2}$	$\bar{7}$	$\bar{3}$	$\bar{8}$	$\bar{4}$
$\bar{6}$	$\bar{0}$	$\bar{6}$	$\bar{3}$	$\bar{0}$	$\bar{6}$	$\bar{3}$	$\bar{0}$	$\bar{6}$	$\bar{3}$
$\bar{7}$	$\bar{0}$	$\bar{7}$	$\bar{5}$	$\bar{3}$	$\bar{1}$	$\bar{8}$	$\bar{6}$	$\bar{4}$	$\bar{2}$
$\bar{8}$	$\bar{0}$	$\bar{8}$	$\bar{7}$	$\bar{6}$	$\bar{5}$	$\bar{4}$	$\bar{3}$	$\bar{2}$	$\bar{1}$

Since every ring can be viewed as a trivially graded ring, $\mathbb{Z}/9\mathbb{Z}$ can be considered as a $\mathbb{Z}$ -graded ring with grading $R_0 = \mathbb{Z}/9\mathbb{Z}$ and $R_n = \{0\}$ for every $n \in \mathbb{Z}, n \neq 0$. Thus, every homogeneous element of $\mathbb{Z}/9\mathbb{Z}$ is either a unit or nilpotent. By

Corollary 2.10, every proper graded ideal in $\mathbb{Z}/9\mathbb{Z}$ is a graded strongly 1-absorbing primary ideal.

The following provides necessary and sufficient conditions for a graded prime ideal of a graded ring to be a graded strongly 1-absorbing primary ideal.

Proposition 2.12. [8] *Let R be a G -graded ring. A graded prime ideal in R is a graded strongly 1-absorbing primary ideal if and only if*

- (1) R is a HUN-ring, or
- (2) R is a graded local ring with a graded maximal ideal M , has exactly one graded prime ideal that is not the graded maximal ideal (i.e., $\text{Grad}(\{0\})$), and every graded M -primary ideal contains M^2 .

Proof. It follows from ([8], Proposition 2.17). □

Next, we extend the results of [9], originally on strongly 1-absorbing primary ideals, to the graded case. A G -graded ring R is called a graded Noetherian if every ascending chain of graded ideals in R terminates. Thus, we obtain the following consequence of Proposition 2.12

Corollary 2.13. *For any graded Noetherian ring R , the following are equivalent :*

- (1) Every graded primary ideal of R is graded strongly 1-absorbing primary.
- (2) R is a HUN ring.

Proof. (1) $\Rightarrow$ (2) Suppose that R is not a HUN ring. Then, by Proposition 2.12, R is a graded local ring with maximal ideal M , and every graded M -primary ideal contains M^2 . By Proposition 1.2, we obtain

$$\begin{aligned} \text{Grad}(M^4) &= \text{Grad}(M) \cap \text{Grad}(M) \cap \text{Grad}(M) \cap \text{Grad}(M) \\ &= \text{Grad}(M) \\ &= \text{Grad}(\text{Grad}(I)), \text{ for some graded } M\text{-primary ideal } I \\ &= \text{Grad}(I) \\ &= M. \end{aligned}$$

Thus, M^4 is an M -primary ideal, and since $M^2 \subseteq M^4$, we get $M^2 = M^4$. Consequently, M^2 is an idempotent ideal. Since R is Noetherian, M^2 is generated by an idempotent element of R . However, because R is a graded local ring, the only idempotent elements in R are 0 and 1. Therefore, we conclude that $M^2 = \{0\}$. As a result, $M^2 \subseteq \text{Grad}(\{0\})$, which implies $M = \text{Grad}(\{0\})$. This means that R is a HUN ring, contradicting our initial assumption. Hence, we conclude that R must be a HUN ring.

(2) $\Rightarrow$ (1) If R is a graded Noetherian ring and a HUN ring, then by Corollary 2.10 every ideal of R is graded strongly 1-absorbing primary. Similarly, every graded primary ideal is also graded strongly 1-absorbing primary. □

Let R be a commutative ring with identity 1_R and let M be an R -module. The idealization of the module M (trivial extension of the ring R by the module M), denote by

$$R(+)M = \{(r, m) \mid r \in R, m \in M\},$$

is a commutative ring with identity $(1_R, 0)$, equipped with componentwise addition and multiplication $(r_1, m_1)(r_2, m_2) = (r_1r_2, r_1m_2 + r_2m_1)$ for all $(r_1, m_1), (r_2, m_2) \in R(+)M$. If $R = \bigoplus_{g \in G} R_g$ be a G -graded ring and $M = \bigoplus_{g \in G} M_g$ be a G -graded R -module, then $R(+)M$ is a G -graded ring with grading $\{R_g(+)M_g\}_{g \in G}$. Suppose that I is an ideal of R and N is a submodule of M . Then $I(+)N$ is an ideal of $R(+)M$ if and only if $IM \subseteq N$. In the context of graded rings, $I(+)N$ is a graded ideal of $R(+)M$ if and only if I is a graded ideal of R and N is a graded submodule of M . Furthermore, the graded radical of a graded $I(+)N$ is defined by $\text{Grad}(I(+)N) = \text{Grad}(I)(+)M$. The following proposition provides a necessary condition for a homogeneous graded ideal $I(+)N$ to be a graded strongly 1-absorbing primary ideal in $R(+)M$.

Proposition 2.14. *Let M be a graded R -module, and let $I(+)N$ be a homogeneous graded ideal of $R(+)M$. If $I(+)N$ is a graded strongly 1-absorbing primary ideal in $R(+)M$, then I is a graded strongly 1-absorbing primary ideal of R .*

Proof. Let $i_1, i_2, i_3 \in h(R)$ be homogeneous non-unit elements such that $i_1i_2i_3 \in I$ and $i_3 \notin \text{Grad}(\{0_R\})$. Then $(i_1, 0)(i_2, 0)(i_3, 0) = (i_1i_2i_3, 0) \in I(+)N$. Since $I(+)N$ is a graded strongly 1-absorbing primary ideal and $(i_3, 0) \notin \text{Grad}(\{0_{R(+)M}\})$, it follows that $(i_1, 0)(i_2, 0) = (i_1i_2, 0) \in I(+)N$, which implies $i_1i_2 \in I$. Hence, I is a graded strongly 1-absorbing primary ideal of R . $\square$

In the following, we introduce a class of ideals that generalizes the concept of graded prime ideals, namely graded 2-absorbing I -ideal, where I is a fixed proper ideal.

Definition 2.15. [10] *Let R be a G -graded ring and let I be a fixed proper ideal of R_e . A proper graded ideal P of R is called a graded 2-absorbing I -ideal if for all $a, b, c \in h(R)$ such that $abc \in P - IP$, then $ab \in P$ or $ac \in P$ or $bc \in P$.*

Theorem 2.16. [10] *If P and Q are non zero graded I -prime ideals of a G -graded ring R , then $P \cap Q$ is a graded 2-absorbing I -ideal.*

Proof. It follows from ([10], Theorem 2.4). $\square$

Proposition 2.17. *Let P be graded strongly 1-absorbing primary ideal of a G -graded ring R and I be a graded ideal. Then $\text{Grad}(P)$ is a graded I -prime ideal of R .*

Proof. Let $a, b \in h(R)$ be arbitrary non-unit homogeneous elements such that $ab \in \text{Grad}(P) - I\text{Grad}(P)$, which means that $ab \in \text{Grad}(P)$ and $ab \notin I\text{Grad}(P)$. Therefore, there exists $n \in \mathbb{N}$ such that $(ab)^n = a^n b^n \in P$. Suppose $n = r + s$ for some $r, s \in \mathbb{N}$. Since P is a graded strongly 1-absorbing primary ideal, it follows that $a^r a^s = a^n \in P$ or $b^n \in \text{Grad}(\{0\}) \subseteq \text{Grad}(P)$. In other words, either $a^n \in P$

or $b^{nm} \in P$ for some $m \in \mathbb{N}$. Consequently, $a \in \text{Grad}(P)$ or $b \in \text{Grad}(P)$. Hence, it is concluded that $\text{Grad}(P)$ is a graded I -prime ideal. $\square$

Proposition 2.18. *Let P and Q be a graded strongly 1-absorbing primary ideal of R . Then $\text{Grad}(PQ)$ is a graded 2-absorbing I -ideal of R .*

Proof. Based on Lemma 1.2, it follows that

$$\text{Grad}(PQ) = \text{Grad}(P \cap Q) = \text{Grad}(P) \cap \text{Grad}(Q).$$

Since P and Q are graded strongly 1-absorbing primary ideals, then by Proposition 2.17, both $\text{Grad}(P)$ and $\text{Grad}(Q)$ are graded I -prime ideals of R . Furthermore, by Theorem 2.16, $\text{Grad}(PQ)$ is a graded 2-absorbing I -ideal. $\square$

It has been shown that the G -graded ring $R \times S$ does not admit any graded strongly 1-absorbing primary ideal. Based on Definition 2.15, we investigate whether the direct product $P \times Q$ forms a graded 2-absorbing $(I \times J)$ -ideal in $R \times S$, where P and Q are graded strongly 1-absorbing primary ideals in R and S , respectively. Let $(a_1, a_2), (b_1, b_2), (c_1, c_2) \in h(R \times S)$ such that

$$(a_1, a_2)(b_1, b_2)(c_1, c_2) = (a_1b_1c_1, a_2b_2c_2) \in (P \times Q) - (I \times J)(P \times Q),$$

which means that $(a_1b_1c_1, a_2b_2c_2) \in P \times Q$ and $(a_1b_1c_1, a_2b_2c_2) \notin (I \times J)(P \times Q)$. Since P and Q are graded strongly 1-absorbing primary ideals in R and S , respectively, it follows that $a_1b_1 \in P$ or $c_1 \in \text{Grad}(\{0_R\})$, and $a_2b_2 \in Q$ or $c_2 \in \text{Grad}(\{0_S\})$. In other words, $(a_1b_1, a_2b_2) \in P \times Q$ or $(c_1, c_2) \in \text{Grad}(\{0_R\}) \times \text{Grad}(\{0_S\})$. Since $\text{Grad}(\{0_R\}) \subseteq \text{Grad}(P)$ and $\text{Grad}(\{0_S\}) \subseteq \text{Grad}(Q)$, we obtain $(a_1b_1, a_2b_2) \in P \times Q$ or $(a_1c_1, a_2c_2) \in \text{Grad}(P) \times \text{Grad}(Q)$ or $(b_1c_1, b_2c_2) \in \text{Grad}(P) \times \text{Grad}(Q)$. This result shows that in general $P \times Q$ does not satisfy the definition of a graded 2-absorbing $(I \times J)$ -ideal. This motivates introducing a new class of graded ideals called graded 2-absorbing I -primary ideals.

Definition 2.19. *Let R be a G -graded ring and I a fixed proper ideal of R_e . A proper graded ideal P of R is called a graded 2-absorbing I -primary ideal if for all $a, b, c \in h(R)$ with $abc \in P - IP$, then $ab \in P$ or $ac \in \text{Grad}(P)$ or $bc \in \text{Grad}(P)$.*

3. CONCLUSIONS

This article establishes several key results. The intersection of any collection of J -graded strongly 1-absorbing primary ideals is again a J -graded 1-absorbing primary ideal. A graded Noetherian ring is a HUN -ring if and only if every graded primary ideal is a graded strongly 1-absorbing primary ideal. A homogeneous graded ideal of the form $I(+)N$ is a graded strongly 1-absorbing primary ideal in the graded idealization $R(+)M$ if and only if I is a graded strongly 1-absorbing primary ideal in R . Furthermore, we introduced generalizations of graded prime ideals, namely graded I -prime ideals and graded 2-absorbing I -ideal, where I is a fixed proper ideal of R . For any graded strongly 1-absorbing primary ideal in a graded ring R , its graded radical is a graded I -prime ideal, and the graded

radical of the product of two graded strongly 1-absorbing primary ideals is a graded 2-absorbing I -ideal. Finally, even if graded rings R and S each contain graded strongly 1-absorbing primary ideals, their direct product $R \times S$ does not necessarily admit such an ideal. However, it was shown that a graded ideal of the form $P \times Q$ in $R \times S$, where P and Q are graded strongly 1-absorbing primary ideals in R and S , respectively, forms a graded 2-absorbing I -primary ideal.

REFERENCES

- [1] F. Farzalipour and P. Ghiasvand, "On the union of graded prime submodules," *Thai J. Math.*, vol. 9, no. 1, pp. 49–55, 2011. [\[1\]](#)
- [2] M. Refai, "Graded radicals and graded prime spectra," *Far East J. Math. Sci.*, pp. 59–73, 2000. [\[2\]](#)
- [3] M. Refai and K. Al-Zoubi, "On graded primary ideals," *Turkish J. Math.*, vol. 28, no. 3, pp. 217–230, 2004. [\[3\]](#)
- [4] K. Al-Zoubi and N. Sharafat, "On graded 2-absorbing primary and graded weakly 2-absorbing primary ideals," *J. Korean Math. Soc.*, vol. 54, no. 2, pp. 675–684, 2007. [\[4\]](#)
- [5] R. Abu-Dawwas and M. Bataineh, "Graded 1-absorbing primary ideals," in *Turkish J. Math. - Studies on Scientific Developments in Geometry, Algebra, and Applied Mathematics*, pp. 1–3, 2022. [\[5\]](#)
- [6] G. Călugăreanu, "Un-rings," *J. Algebra Appl.*, vol. 15, 2015. [\[6\]](#)
- [7] A. S. Alshehry, R. Abu-Dawwas, and M. Al-Rashdan, "Some notes on graded weakly 1-absorbing primary ideals," *Demonstratio Math.*, vol. 56, no. 1, pp. 42–52, 2023. [\[7\]](#)
- [8] R. Abu-Dawwas, "On graded strongly 1-absorbing primary ideals," *Khayyam J. Math.*, vol. 8, no. 1, pp. 42–52, 2022. [\[8\]](#)
- [9] F. A. A. Almahdi, E. M. Bouba, and A. N. A. Koam, "On strongly 1-absorbing primary ideals of commutative rings," *Bull. Korean Math. Soc.*, vol. 57, no. 5, pp. 1250–1213, 2020. [\[9\]](#)
- [10] I. Akray, A. K. Jabbar, and S. A. Othman, "Graded n-absorbing i-ideals," *Palestine J. Math.*, vol. 13, no. 1, 2024. [\[10\]](#)
- [11] A. Badawi and E. Y. Celikel, "On 1-absorbing primary ideals of commutative rings," *J. Algebra Appl.*, vol. 19, no. 6, 2020. [\[11\]](#)

Secret Sharing Schemes from Two Families of Cyclic Codes Using Massey's Construction

Syahrul Zada^{1*}, Karyati²

^{1,2}Universitas Negeri Yogyakarta, Indonesia

¹syahrulzada.2020@student.uny.ac.id, ²karyati@uny.ac.id

Abstract. The objectives of this study include proving that $\tilde{C}_{(q,m,\delta_2)}$ and $\tilde{C}_{(q,m,\delta_3)}$, which are Primitive BCH codes, with $m \geq 5$ are minimal codes, and presenting specific examples of secret-sharing schemes based on dual of these codes. To prove that $\tilde{C}_{(q,m,\delta_2)}$ and $\tilde{C}_{(q,m,\delta_3)}$ with $m \geq 5$ are minimal codes, the criterion $\frac{W_{min}}{W_{max}} > \frac{q-1}{q}$ used, where W_{min} and W_{max} are the minimum weight and maximum weight, respectively. Data on the minimum weight and maximum weight of $\tilde{C}_{(q,m,\delta_2)}$ and $\tilde{C}_{(q,m,\delta_3)}$ are obtained from previous research. To give an example of secret-sharing scheme construction based on these codes, the construction method to be used is Massey construction. This research successfully proves that $\tilde{C}_{(q,m,\delta_2)}$ and $\tilde{C}_{(q,m,\delta_3)}$ with $m \geq 5$ are minimal codes. In addition, this research also successfully presents an example of secret-sharing scheme construction based on these codes using Massey's construction.

Keywords: dual codes, Massey's construction, minimal codes, primitive BCH codes, secret-sharing.

1. INTRODUCTION

The *secret sharing* scheme is one of the protocols in cryptography that aims to share secret data with several parties, where the parties receiving the secret sharing must work together to access the secret data. The idea of secret sharing was first pioneered by Shamir [1]. The scheme he introduced is called the (k, n) threshold scheme

The (k, n) threshold scheme uses polynomial interpolation in the sharing and recovery of *secret*. Besides using polynomial interpolation, secret sharing schemes

*Corresponding author

2020 Mathematics Subject Classification: 55S30

Received: 01-07-2025, accepted: 21-08-2025.

can be built using other mathematical objects. The idea of replacing polynomial interpolation with other algorithms was first proposed by McEliece & Sarwate [2]. They replaced polynomial interpolation on the (k, n) threshold scheme with a Reed-Solomon code encoding and decoding algorithm.

Apart from using Reed-Solomon codes, the *secret-sharing* scheme can be constructed using linear codes in general. In the existing literature, there are two approaches to construct textitsecret sharing schemes based on linear codes. The first approach was developed in 1989 by Brickell [3]. While the second approach was developed by Massey [4].

In the second construction or commonly called Massey construction, the minimal *codeword* is needed to determine the minimal access set (the smallest set of participants that can select the secret). Therefore, the minimal *codeword* in a linear code needs to be found. Finding all minimal codewords of a linear code is quite a difficult problem as it requires testing q^k codewords of a linear code. The problem is one form of the *covering problem*.

One way to simplify the *covering problem* on linear codes is to use the Ashikhmin-Barg [5] criterion. If a linear code satisfies this criterion, it is called a minimal linear code. A minimal linear code is a type of linear code that can produce a *secret sharing* scheme with an interesting access structure [6].

An example of a minimal linear code has been produced by Ding, Fan, and Zhou [7], namely the $\tilde{C}_{(q,m,\delta_2)}$ code and the $\tilde{C}_{(q,m,\delta_3)}$ code with $m \geq 5$ which are Primitive BCH codes with designed distances δ_2 and δ_3 . Let $m > 1$ be a positive integer, and let $n = q^m - 1$. Suppose α is the generator of $\mathbb{F}_{q^m}^*$, which is the multiplicative group of $\mathbb{F}_{q^m}$. For every i with $0 \leq i \leq q^m - 2$, let $m_i(x)$ denote the minimal polynomial of α^i over $\mathbb{F}_{q^m}$. For every $2 \leq \delta < n$, define

$$g_{q,m,\delta}(x) = \text{KPK}(m_1(x), m_2(x), \dots, m_{\delta-1}(x)),$$

where KPK denotes the least common multiple. In addition, also define

$$\tilde{g}_{q,m,\delta}(x) = (x - 1)g_{q,m,\delta}(x).$$

Let $C_{(q,m,\delta)}$ and $\tilde{C}_{(q,m,\delta)}$ denote cyclic codes of length n with generator polynomials $g_{(q,m,\delta)}(x)$ and $\tilde{g}_{(q,m,\delta)}(x)$, respectively. The set $C_{(q,m,\delta)}$ is a primitive BCH code with designed distance δ , and $\tilde{C}_{(q,m,\delta)}$ is a primitive BCH code with designed distance δ .

Actually, Ding et al. discuss the dimensions and weights of two families of BCH codes. However, in the last part of their article, Ding et al. the code $\tilde{C}_{(q,m,\delta_2)}$ and the code $\tilde{C}_{(q,m,\delta_3)}$ satisfy the Ashikhmin-Barg criterion when $m \geq 5$ so they are minimal codes. Unfortunately, the statement has not been accompanied by a proof. Thus, in this paper, the proof that these 2 codes are minimal codes will be presented.

2. Weight Distribution and Parameter of Code

Before presenting the characterization results of primitive BCH codes with δ_2 and δ_3 designed distances, the following two theorems are first presented. These two theorems have been proved in [7]. The following theorem provides information about the weight distribution of the code $\tilde{C}_{(q,m,\delta_2)}$.

Theorem 2.1. [7] *The code $\tilde{C}_{(q,m,\delta_2)}$ has parameters $[n, \tilde{k}, \tilde{d}]$, where $\tilde{d} \geq \delta_2 + 1$ and*

$$\tilde{k} = \begin{cases} 2m & \text{for odd } m, \\ \frac{3m}{2} & \text{for even } m. \end{cases}$$

When q is an odd prime, $\tilde{d} = \delta_2 + 1$ and $\tilde{C}_{(q,m,\delta_2)}$ is three-weight code with the weight distribution of Table 1 for odd m and Table 2 for even m .

TABLE 1. Weight distribution of $\tilde{C}_{(q,m,\delta_2)}$ for odd m

Weight w	Number of codeword A_w
0	1
$(q-1)q^{m-1} - q^{(m-1)/2}$	$(q-1)(q^{m-1})(q^{m-1} + q^{(m-1)/2})/2$
$(q-1)q^{m-1}$	$(q^m - 1)(q^{m-1} + 1)$
$(q-1)q^{m-1} + q^{(m-2)/2}$	$(q-1)(q^{m-1})(q^{m-1} - q^{(m-1)/2})/2$

TABLE 2. Weight distribution of $\tilde{C}_{(q,m,\delta_2)}$ for even m

Weight w	Number of codeword A_w
0	1
$(q-1)q^{m-1} - q^{(m-1)/2}$	$(q-1)(q^{(3m-2)/2} - q^{(m-2)/2})$
$(q-1)q^{m-1}$	$q^m - 1$
$(q-1)(q^{m-1} + q^{(m-2)/2})$	$q^{(m-2)/2}(q^m - q^{(m+2)/2} + q - 1)$

The following theorem informs about the weight distribution of the code $\tilde{C}_{(q,m,\delta_3)}$.

Theorem 2.2. [7] *Let $m \geq 4$. The code $\tilde{C}_{(q,m,\delta_3)}$ has parameters $[n, \tilde{k}, \tilde{d}]$, where $\tilde{d} \geq \delta_3 + 1 = (q-1)q^m - 1 - q^{\lfloor (m+1)/2 \rfloor}$ and*

$$\tilde{k} = \begin{cases} 2m & \text{for odd } m, \\ \frac{5m}{2} & \text{for even } m. \end{cases}$$

When q is an odd prime and $m \geq 4$ is even, the code $\tilde{C}_{(q,m,\delta_3)}$ has minimum distance $\tilde{d} = \delta_3 + 1$ and its weight distribution is given in Table 3. When q is an odd prime and $m \geq 5$ is odd, the code $\tilde{C}_{(q,m,\delta_3)}$ has minimum distance $\tilde{d} = \delta_3 + 1$ and its weight distribution is given in Table 4.

TABLE 3. The weight distribution of $\tilde{C}_{(q,m,\delta_3)}$ for even m and odd q

Weight w	Number of codeword A_w
0	1
$(q-1)q^{m-1} - q^{m/2}$	$(q^m - 1) \left((q^2 - 1) \left(q^{(3m-6)/2} + q^{m-2} \right) + 2 \left(q^{(m-2)/2} - 1 \right) \left(q^{m-3} + q^{(m-4)/2} \right) \right) / 2(q+1)$
$(q-1) \left(q^{m-1} - q^{(m-2)/2} \right)$	$q \left(q^{m/2} + 1 \right) (q^m - 1) \left(q^{m-1} + (q-1)q^{(m-2)/2} \right) / 2(q+1)$
$(q-1)q^{m-1} - q^{(m-2)/2}$	$\left(q^{m+1} - 2q^m + q \right) \left(q^{m/2} - 1 \right) \left(q^{m-1} + q^{(m-2)/2} \right) / 2$
$(q-1)q^{m-1}$	$(q^m - 1) \left(1 + q^{(3m-2)/2} - q^{(3m-4)/2} + 2q^{(3m-6)/2} - q^{m-2} \right)$
$(q-1)q^{m-1} + q^{(m-2)/2}$	$q \left(q^{m/2} + 1 \right) (q^m - 1) (q-1) \left(q^{m-1} - q^{(m-2)/2} \right) / 2(q+1)$
$(q-1) \left(q^{m-1} + q^{(m-2)/2} \right)$	$\left(q^{m+1} - 2q^m + q \right) \left(q^{m/2} - 1 \right) \left(q^{m-1} - (q-1)q^{(m-2)/2} \right) / 2(q-1)$
$(q-1)q^{m-1} + q^{m/2}$	$q^{(m-2)/2} (q^m - 1) (q-1) \left(q^{m-2} - q^{(m-2)/2} \right) / 2$
$(q-1) \left(q^{m-1} + q^{m/2} \right)$	$\left(q^{(m-2)/2} - 1 \right) (q^m - 1) \left(q^{m-3} - (q-1)q^{(m-4)/2} \right) / (q^2 - 1)$

TABLE 4. The weight distribution of $\tilde{C}_{(q,m,\delta_3)}$ for odd m and odd q

Weight w	Number of codeword A_w
0	1
$(q-1)q^{m-1} - q^{(m+1)/2}$	$(q^m - 1) \left(q^{m-3} + q^{(m-3)/2} \right) \left(q^{m-1} - 1 \right) / 2(q+1)$
$(q-1) \left(q^{m-1} - q^{(m-1)/2} \right)$	$(q^m - 1) \left(q^{m-1} + q^{(m-1)/2} \right) \left(q^{m-2} + (q-1)q^{(m-3)/2} \right) / 2$
$(q-1)q^{m-1} - q^{(m-1)/2}$	$(q^m - 1) \left(q^{m-2} + q^{(m-3)/2} \right) \left(q^{m+3} - q^{m+2} - q^{m-1} - q^{(m+3)/2} + q^{(m-1)/2} + q^3 \right) / 2(q+1)$
$(q-1)q^{m-1}$	$(q^m - 1) \left(1 + (q^2 - q + 1) q^{m-3} + (q-1)q^{2m-4} + (q-2)q^{2m-2} + q^{2m-1} \right)$
$(q-1)q^{m-1} + q^{(m-1)/2}$	$(q^m - 1) \left(q^{m-2} - q^{(m-3)/2} \right) \left(q^{m+3} - q^{m+2} - q^{m-1} + q^{(m+3)/2} - q^{(m-1)/2} + q^3 \right) / 2(q+1)$
$(q-1) \left(q^{m-1} + q^{(m-1)/2} \right)$	$(q^m - 1) \left(q^{m-1} - q^{(m-1)/2} \right) \left(q^{m-2} - (q-1)q^{(m-3)/2} \right) / 2$
$(q-1)q^{m-1} + q^{(m+1)/2}$	$(q^m - 1) \left(q^{m-3} - q^{(m-3)/2} \right) \left(q^{m-1} - 1 \right) / 2(q+1)$

3. Minimal Codeword

The following concepts of *support* and *covering* are the origin of minimal linear codes. The definition of *support* is explained as follows.

Definition 3.1. [6] The **support** of $\mathbf{c} \in \mathbb{F}_q^n$ is defined by

$$\text{supp}(\mathbf{c}) = \{0 \leq i \leq n-1 | c_i \neq 0\}.$$

Definition 3.2. [6] A vector $u \in \mathbb{F}_q^n$ covers a vector $v \in \mathbb{F}_q^n$ if $\text{supp}(v)$ is subset of $\text{supp}(u)$.

Minimal code is defined as follows.

Definition 3.3. [6] A nonzero codeword $\mathbf{u}$ in a linear code C is minimal if $\mathbf{u}$ covers only scalar multiples of $\mathbf{u}$, but no other nonzero codewords in C . A linear code C is minimal if every nonzero codeword in C is minimal.

The following lemma is often used in Linear Code Characterization. This lemma will also be used in this paper

Lemma 3.4 (Ashikhmin-Barg). [6] A $[n, k, d]_q$ linear code C is minimal if

$$\frac{W_{\min}}{W_{\max}} > \frac{q-1}{q}. \quad (1)$$

where $W_{\min}$ and $W_{\max}$ denote the minimum and maximum nonzero Hamming weights of C respectively.

Using Lemma 3.4, many families of minimal linear codes with $\frac{W_{min}}{W_{max}} > \frac{q-1}{q}$ have been found, such as those by Carlet, Ding and Yuan [8]. However, most of these codes have a limited number of weights. For example, [9] introduced a type of cyclic code that has three weights, where each *codeword* has one of the three different weights. The code satisfies this lemma.

4. Massey's Construction

Consider the linear code $C[n, k, d]_q$. Suppose $G = [\mathbf{g}_0, \mathbf{g}_1, \dots, \mathbf{g}_{n-1}]$ the generator matrix of C . The secret S is a member of $\mathbb{F}_q$. *Share* can be determined by the following procedure. Randomly select the vector $\mathbf{u} = (u_0, u_1, \dots, u_{k-1}) \in \mathbb{F}_q^k$ such that $S = \mathbf{u}\mathbf{g}_0$. Then, the vector $\mathbf{s}$ can be calculated by

$$\mathbf{s} = (S, s_1, \dots, s_{n-1}) = \mathbf{u}G.$$

share for each participant P_i is s_i for all $1 \leq i \leq n-1$.

Assume m people collect their respective *share* $\{s_{i_1}, s_{i_2}, \dots, s_{i_m}\}$ with $1 \leq m \leq n-1$. Then, the secret $S = s_0 + \mathbf{u}\mathbf{g}_0$ can be determined if and only if $\mathbf{g}_0$ is a linear combination of $\mathbf{g}_1, \mathbf{g}_2, \dots, \mathbf{g}_{n-1}$. Thus, resulting in the following proposition.

Proposition 4.1. *Let G be a generator matrix of an $[n, k, d]_q$ linear code C . In the secret-sharing based on C with respect to the second construction, a set of share $\{s_{i_1}, s_{i_2}, \dots, s_{i_m}\}$ with $1 \leq i_1 < i_2 < \dots < i_m \leq n-1$ and $1 \leq m \leq n-1$, determines the secret if and only if there is a codeword*

$$\mathbf{c} = (1, 0, \dots, 0, c_{i_1}, 0, \dots, 0, c_{i_m}, 0, \dots, 0)$$

in the dual code $C^\perp$, where $c_{i_j} \neq 0$ for at least one j .

5. The proof of $\tilde{C}_{(q,m,\delta_2)}$ with $m \geq 5$ is a minimal code

In this section, we will present the results of proving that Primitive BCH codes with designed distances δ_2 and δ_3 with $m \geq 5$ satisfy Lemma 3.4, so they are minimal codes.

The first result of this research is presented in the following theorem along with its proof. The following theorem states that the Primitive BCH code with Designed Distance δ_2 satisfies Lemma 3.4.

Theorem 5.1. *Let $\delta_2 = (q-1)q^{m-1} - 1 - q^{(m-1)/2}$,*

- (1) *If $m \geq 5$ and m is odd, $\tilde{C}_{(q,m,\delta_2)}$ is minimal code.*
- (2) *If $m \geq 6$ and m is even, $\tilde{C}_{(q,m,\delta_2)}$ is minimal code.*

Proof. (a) From the Table 1, we know that $\tilde{C}_{(q,m,\delta_2)}$ with odd m has nonzero minimum weight $W_{min} = (q-1)q^{m-1} - q^{(m-1)/2}$ and maximum weight $W_{max} = (q-1)q^{m-1} + q^{(m-2)/2}$.

Using Lemma 3.4, we will proof that

$$\frac{(q-1)q^{m-1} - q^{(m-1)/2}}{(q-1)q^{m-1} + q^{(m-1)/2}} > \frac{q-1}{q}.$$

Suppose $A = (q-1)q^{m-1}$ and $B = q^{(m-1)/2}$. So that the inequality becomes

$$\begin{aligned} \frac{A-B}{A+B} &> \frac{q-1}{q} \\ (A-B)q &> (A+B)(q-1) \\ Aq - Bq &> Aq - A + Bq - B \\ A &> 2Bq - B \\ A &> B(2q-1) \end{aligned}$$

Substitute $A = (q-1)q^{m-1}$ and $B = q^{(m-1)/2}$ back into the inequality above.

$$\begin{aligned} (q-1)q^{m-1} &> q^{(m-1)/2}(2q-1). \\ (q-1)q^{(m-1)/2} &> 2q-1. \end{aligned}$$

Since q is an odd prime number and $m > 5$ is odd, $q^{(m-1)/2}$ is a large positive number, making the left side larger than the right side. Thus, it is proved that for q odd primes and $m > 5$ odd,

$$\frac{(q-1)q^{m-1} - q^{(m-1)/2}}{(q-1)q^{m-1} + q^{(m-1)/2}} > \frac{q-1}{q}.$$

(b) From Table 2, it can be seen that $\tilde{C}_{(q,m,\delta_2)}$ with even m have a nonzero minimum weight ($W_{min} = (q-1)q^{m-1} - q^{(m-1)/2}$) and maximum weight ($W_{max} = (q-1)(q^{m-1} + q^{(m-2)/2})$).

Using Lemma 3.4, we will proof that

$$\frac{(q-1)q^{m-1} - q^{(m-2)/2}}{(q-1)(q^{m-1} + q^{(m-2)/2})} > \frac{q-1}{q}.$$

with algebraic manipulation the above inequality becomes

$$\frac{(q-1)q^{m-1} - q^{(m-2)/2}}{q^{m-1} + q^{(m-2)/2}} > \frac{(q-1)^2}{q}.$$

Suppose $A = q^{m-1}$ and $B = q^{(m-2)/2}$. So that the above inequality becomes

$$\begin{aligned}
\frac{(q-1)A-B}{A+B} &> \frac{(q-1)^2}{q} \\
q\{(q-1)A-B\} &> (q-1)^2(A+B) \\
A(q^2-q) - Bq &> (q^2-2q+1)(A+B) \\
Aq^2 - Aq - Bq &> Aq^2 - 2Aq + A + Bq^2 - 2Bq + B \\
Aq + Bq &> A + Bq^2 + B \\
Aq - A &> Bq^2 + B - Bq \\
(q-1)A &> (q^2-q+1)B
\end{aligned}$$

substitute $A = q^{m-1}$ and $B = q^{(m-2)/2}$ back into the inequality above.

$$\begin{aligned}
(q-1)q^{m-1} &> (q^2-q+1)q^{(m-2)/2} \\
(q-1)q^{m/2} &> (q^2-q+1) \\
(q-1)q^{m/2} - 1 &> q^2 - q \\
(q-1)q^{m/2} - (q-1)\frac{1}{q-1} &> (q-1)q \\
q^{m/2} - \frac{1}{q-1} &> q
\end{aligned}$$

since q is an odd prime number then $0 < \frac{1}{q-1} \leq 1$, and since $m \geq 5$ then $q^{m/2} > q$, so the left side will always be greater than the right side. So it is proven that

$$\frac{(q-1)q^{m-1} - q^{(m-2)/2}}{(q-1)(q^{m-1} + q^{(m-2)/2})} > \frac{q-1}{q}.$$

□

The next result of this research is presented in the following theorem and its proof.

6. Proof of $\tilde{C}_{(q,m,\delta_3)}$ with $m \geq 5$ is a minimal code

The following theorem states that the Primitive BCH code with Designed Distance δ_3 satisfies Lemma 3.4.

Theorem 6.1. Let $\delta_2 = (q-1)q^{m-1} - 1 - q^{(m+1)/2}$,

- (1) If $m \geq 6$ and m is even, then $\tilde{C}_{(q,m,\delta_3)}$ is minimal code.
- (2) If $m \geq 5$ and m is odd, then $\tilde{C}_{(q,m,\delta_3)}$ is minimal code.

Proof. (a) From the table 3, can be seen that for even m the code $\tilde{C}_{(q,m,\delta_3)}$ have nonzero minimal weight $(q-1)q^{m-1} - q^{m/2}$ and maximum weight $(q-1)(q^{m-1} + q^{m/2})$.

Using Lemma 3.4, we will show that

$$\frac{(q-1)q^{m-1} - q^{m/2}}{(q-1)(q^{m-1} + q^{m/2})} > \frac{q-1}{q}$$

multiply both side by $(q-1)$ so that

$$\begin{aligned} \frac{(q-1)q^{m-1} - q^{m/2}}{(q-1)(q^{m-1} + q^{m/2})} \cdot (q-1) &> \frac{q-1}{q} \cdot (q-1) \\ \frac{(q-1)q^{m-1} - q^{m/2}}{(q^{m-1} + q^{m/2})} &> \frac{(q-1)^2}{q} \end{aligned}$$

suppose $A = q^{m-1}$ and $B = q^{m/2}$. So that the inequality becomes

$$\begin{aligned} \frac{(q-1)A - B}{A + B} &> \frac{(q-1)^2}{q} \\ q\{(q-1)A - B\} &> (q-1)^2(A + B) \\ A(q^2 - q) - Bq &> (q^2 - 2q + 1)(A + B) \\ Aq^2 - Aq - Bq &> Aq^2 - 2Aq + A + Bq^2 - 2Bq + B \\ Aq + Bq &> A + Bq^2 + B \\ Aq - A &> Bq^2 - Bq + B \\ (q-1)A &> (q^2 - q + 1)B \end{aligned}$$

sustitute $A = q^{m-1}$ and $B = q^{m/2}$ back into the inequality above

$$\begin{aligned} (q-1)q^{m-1} &> (q^2 - q + 1)q^{m/2} \\ (q-1)q^{(m-2)/2} &> (q^2 - q + 1) \\ (q-1)q^{(m-2)/2} - 1 &> q^2 - q \\ (q-1)q^{(m-2)/2} - (q-1)\frac{1}{q-1} &> (q-1)q \\ q^{(m-2)/2} - \frac{1}{q-1} &> q \\ q^{(m-2)/2} &> q + \frac{1}{q-1} \end{aligned}$$

Since q is an odd prime number and $m > 5$ so the left side will always greater than right side. So it is proven that for q odd prime number and even $m > 5$

$$\frac{(q-1)q^{m-1} - q^{m/2}}{(q-1)(q^{m-1} + q^{m/2})} > \frac{q-1}{q}$$

(b) From the table 4, we know that for odd m , $\tilde{C}_{(q,m,\delta_3)}$ have nonzero minimal weight $(q-1)q^{m-1} - q^{(m+1)/2}$ and maximum weight $(q-1)q^{m-1} + q^{(m-1)/2}$.

Using Lemma 3.4, we will prove that

$$\frac{(q-1)q^{m-1} - q^{(m+1)/2}}{(q-1)q^{m-1} + q^{(m+1)/2}} > \frac{q-1}{q}$$

let $A = (q-1)q^{m-1}$ and $B = q^{(m+1)/2}$, then

$$\begin{aligned}\frac{A-B}{A+B} &> \frac{q-1}{q} \\ (A-B)q &> (A+B)(q-1) \\ Aq - Bq &> Aq - A + Bq - B \\ A &> 2Bq - B \\ A &> B(2q-1)\end{aligned}$$

Substitute $A = (q-1)q^{m-1}$ and $B = q^{(m+1)/2}$ back to the inequality above

$$\begin{aligned}(q-1)q^{m-1} &> q^{(m+1)/2}(2q-1) \\ (q-1)q^{(m^2-1)/2} &> 2q-1\end{aligned}$$

Because (q) is odd prime number and $(m \geq 5)$, $(q^{(m^2-1)/2})$ is a large positive number, making the left side larger than the right side. Thus, it is proved that for q odd primes, $m \geq 5$ and odd m ,

$$\frac{(q-1)q^{m-1} - q^{(m+1)/2}}{(q-1)q^{m-1} + q^{(m+1)/2}} > \frac{q-1}{q}$$

□

7. Example of Secret Sharing Schemes Based On Dual Code of $\tilde{C}_{(q,m,\delta_2)}$ with $m \geq 5$ using Massey's Construction

The next goal of this paper is to provide an example of a secret sharing scheme based on the dual code of $\tilde{C}_{(q,m,\delta_2)}$ with $m \geq 5$ using Massey's Construction.

Example 7.1. Let $q = 2$, $m = 5$, then $n = q^m - 1 = 31$ and designed distance $\delta_2 = (q-1)q^{m-1} - 1 - q^{(m+1)/2} = 11$. With the commonly known technique to construct the generator matrix of cyclic codes, the generator matrix of $\tilde{C}_{(2,5,11)}$ is obtained, that is

$$\tilde{G} = \begin{bmatrix} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 1 & 1 & 1 & 1 & 0 & 1 & 1 & 0 & 1 & 0 & 1 & 1 & 0 & 0 & 1 & 0 & 1 & 1 & 1 & 1\end{bmatrix}$$

Based on matrix $\tilde{G}$ and Proposition 4.1, there is no dictator participant in the example textitsecret-sharing scheme based on the code dual of code $\tilde{C}_{(q,m,\delta_2)}$. The participant P_i with $1 \leq i \leq n-1$ is in $(q-1)q^{k-2} = (2-1)2^{10-2} = 2^8 = 256$ of $q^{k-1} = 2^{10-1} = 512$ minimal access set.

By using the common method, the parity check matrix of $\tilde{C}_{(2,5,11)}$ is obtained, that is

[illegible]

Suppose the secret is 0. The Massey's construction will be applied to the dual code of $\tilde{C}_{(q,m,\delta_2)}$ to construct the secret sharing scheme. A vector $\mathbf{u} \in \mathbb{F}_2^{21}$ needs to be chosen such that $S = \mathbf{u}\tilde{\mathbf{h}}_0$, where $\tilde{\mathbf{h}}_0$ is the first column of the matrix $\tilde{H}$. Misalkan $\mathbf{u} = (0, 1, 0, 1, 0, 1, 0, 1, 0, 1, 0, 1, 0, 1, 0, 1, 0, 1, 0)$. Then the following share can be obtained:

$$\begin{aligned} \mathbf{s} &= \mathbf{u}\tilde{H} \\ &= (S, 1, 1, 1, 0, 1, 0, 1, 0, 1, 1, 0, 1, 0, 1, 0, 1, 0, 1, 1, 0, 1, 1, 1, 1, 1, 1, 0). \end{aligned}$$

Suppose there are 30 participants $P_1, P_2, \dots, P_{30}$. Each participant P_i gets s_i as its share. Now, find the set of participants who can reconstruct the secret. That is, the set of participants such that $\tilde{\mathbf{h}}_0$ is a linear combination of the column of $\tilde{H}$ corresponding to the participants in the set. Here are 3 of $q^{k-1} = 2^{10-1} = 512$ sets of participants that can produce secrets: $\{P_1, P_2, P_3, P_4, P_5, P_6, P_8, P_9, P_{11}, P_{13}, P_{14}, P_{17}, P_{19}, P_{20}, P_{21}\}$, $\{P_1, P_2, P_3, P_4, P_5, P_6, P_8, P_{10}, P_{12}, P_{10}, P_{15}, P_{18}, P_{19}, P_{21}, P_{22}, P_{23}, P_{26}, P_{28}, P_{29}, P_{30}\}$, and $\{P_1, P_2, P_3, P_4, P_7, P_{10}, P_{16}, P_{17}, P_{18}, P_{20}, P_{21}, P_{22}, P_{24}, P_{25}, P_{26}\}$

Based on Proposition 4.1, these participant sets can produce secret. Other minimal access sets can be computed using programming.

Suppose the group of participants recovering share is $\{P_1, P_2, P_3, P_4, P_5, P_6, P_8, P_9, P_{11}, P_{13}, P_{14}, P_{17}, P_{19}, P_{20}, P_{21}\}$, can be seen that $\tilde{\mathbf{h}}_0 = \{1 \cdot \tilde{\mathbf{h}}_1 + 1 \cdot \tilde{\mathbf{h}}_2 + 1 \cdot \tilde{\mathbf{h}}_3 + 1 \cdot \tilde{\mathbf{h}}_4 + 1 \cdot \tilde{\mathbf{h}}_5 + 1 \cdot \tilde{\mathbf{h}}_6 + 1 \cdot \tilde{\mathbf{h}}_8 + 1 \cdot \tilde{\mathbf{h}}_9 + 1 \cdot \tilde{\mathbf{h}}_{11} + 1 \cdot \tilde{\mathbf{h}}_{13} + 1 \cdot \tilde{\mathbf{h}}_{14} + 1 \cdot \tilde{\mathbf{h}}_{17} + 1 \cdot \tilde{\mathbf{h}}_{19} + 1 \cdot \tilde{\mathbf{h}}_{20} + 1 \cdot \tilde{\mathbf{h}}_{21}\}$

So that we get share $S = \mathbf{uh}_0 = 1 \cdot 1 + 1 \cdot 1 + 1 \cdot 1 + 1 \cdot 0 + 1 \cdot 1 + 1 \cdot 0 + 1 \cdot 0 + 1 \cdot 1 + 1 \cdot 0 + 1 \cdot 0 + 1 \cdot 1 + 1 \cdot 1 = 0$.

8. Example of Secret Sharing Schemes Based On Dual Code of $\tilde{C}_{(q,m,\delta_3)}$ with $m \geq 5$ using Massey's Construction

The next goal of this paper is to give an example of a secret sharing scheme based on the dual code of $\tilde{C}_{(q,m,\delta_3)}$ with $m \geq 5$ with Massey construction.

Example 8.1. Let $q = 2$, $m = 5$, then $n = 2^m - 1 = 2^5 - 1 = 31$ and designed distance $\delta_3 = (q - 1)q^{m-1} - 1 - q^{(m+1)/2} = 7$.

By constructing the generator matrix of a commonly known cyclic code, the generator matrix of the code $\tilde{C}_{(2,5,7)}$ is obtained, that is

[illegible]

Based on matrix $\tilde{G}$ and Proposition 4.1, there is no dictator participant in this example. The participant P_i with $1 \leq i \leq n-1$ is in $(q-1)q^{k-2} = (2-1)2^{15-2} = 2^{13} = 8192$ of $q^{k-1} = 2^{15-1} = 16384$ minimal access set.

In a common way, one can also obtain the parity check matrix of $\widetilde{C}_{(2,5,7)}$, as follows

[illegible]

In this step, the secret sharing scheme will be built. Since in constructing the secret sharing scheme the dual code of $\tilde{C}(q, m, \delta_3)$ will be used, the parity check matrix of the code will be used as the generator matrix. Suppose the secret is 1. The Massey's construction will be applied to this code to build the secret sharing scheme. The vector $\mathbf{u} \in \mathbb{F}_2^{16}$ needs to be chosen such that $S = \mathbf{u}\tilde{\mathbf{h}}_0$, where $\tilde{\mathbf{h}}_0$ is the first column of the matrix $\tilde{H}$. Let $\mathbf{u} = (1, 1, 1, 1, 0, 0, 0, 0, 1, 1, 1, 1, 0, 0, 0, 0)$. Then the following share can be obtained:

$$\begin{aligned} \mathbf{s} &= \mathbf{u}\tilde{H} \\ &= (S, 0, 1, 0, 1, 1, 0, 0, 0, 1, 1, 0, 0, 1, 1, 1, 0, 0, 1, 0, 1, 1, 1, 1, 1, 0, 0, 0, 0). \end{aligned}$$

Suppose there are 30 participants $P_1, P_2, \dots, P_{30}$. Each participant P_i gets s_i as its share. Now find the set of participants who can reconstruct the secret. That is, the set of participants such that $\tilde{\mathbf{h}}_0$ is a linear combination of the columns of $\tilde{H}$ corresponding to the participants in the set. Here are 3 of $q^{k-1} = 2^{15-1} = 16384$ sets of participants that can produce secrets: $\{P_4, P_5, P_6, P_7, P_{12}, P_{15}, P_{16}\}$, $\{P_4, P_5, P_6, P_7, P_{12}, P_{14}, P_{15}, P_{16}, P_{18}, P_{19}, P_{20}, P_{21}, P_{26}, P_{29}, P_{30}\}$, and $\{P_8, P_{16}, P_{18}, P_{22}, P_{23}, P_{25}, P_{26}\}$

Example of the set of participants that can produce the secret are obtained from the codeword on $\tilde{C}_{(2,5,7)}$ whose first component is 1. While the search for the codeword can be done with programming.

Suppose the group of participants recovering share is $\{P_4, P_5, P_6, P_7, P_{12}, P_{15}, P_{16}\}$, can be seen that $\tilde{\mathbf{h}}_0 = 1 \cdot \tilde{\mathbf{h}}_4 + 1 \cdot \tilde{\mathbf{h}}_5 + 1 \cdot \tilde{\mathbf{h}}_6 + 1 \cdot \tilde{\mathbf{h}}_7 + 1 \cdot \tilde{\mathbf{h}}_{12} + 1 \cdot \tilde{\mathbf{h}}_{15} + 1 \cdot \tilde{\mathbf{h}}_{16} + 1 \cdot$. So, we get share $S = \mathbf{u}\tilde{\mathbf{h}}_0 = 1 \cdot 1 + 1 \cdot 1 + 1 \cdot 0 + 1 \cdot 0 + 1 \cdot 0 + 1 \cdot 1 + 1 \cdot 0 + 1 \cdot 0 = 1$.

9. CLOSING

From this study it is proved that $\tilde{C}_{(q,m,\delta_2)}$ and $\tilde{C}_{(q,m,\delta_3)}$ with $m \geq 5$ are minimal codes. In addition, this research also successfully present an example of the construction of a secret sharing scheme based on the dual code of the code using Massey's construction. Nevertheless, there are still open problem related to linear code based on secret sharing schemes. One of them is to prove the minimality of the codes $\tilde{C}_{(q,m,\delta_2)}$ and $\tilde{C}_{(q,m,\delta_3)}$ when $m \leq 5$. Besides using the Ashikhmin-Barg criterion, other criteria can also be used as in [10] and [11].

REFERENCES

- [1] A. Shamir, "How to share a secret," *Communications of the ACM*, vol. 22, no. 11, pp. 612–613, 1979.
- [2] R. J. McEliece and D. V. Sarwate, "On sharing secrets and reed-solomon codes," *Commun. ACM*, vol. 24, pp. 583–584, 1981.
- [3] E. Brickell, *Some Ideal Secret Sharing Schemes in : J. Quisquater Jean-Jacques and Vandewalle (Ed.)*. Springer Berlin Heidelberg, 1990.

- [4] J. L. Massey, "Minimal codewords and secret sharing," in *Proceedings of the 6th joint Swedish-Russian international workshop on information theory*, pp. 276–279, 1993. [↗](#)
- [5] A. Ashikhmin and A. Barg, "Minimal vectors in linear codes," *IEEE Transactions on Information Theory*, vol. 44, no. 5, pp. 2010–2017, 1998. [↗](#)
- [6] W. C. Huffman, J.-L. Kim, and P. Solé, *Concise Encyclopedia of Coding Theory*. Chapman and Hall/CRC, 2021. [↗](#)
- [7] C. Ding, C. Fan, and Z. Zhou, "The dimension and minimum distance of two classes of primitive bch codes," *Finite Fields Their Appl*, vol. 45, pp. 237–263, 2016. [↗](#)
- [8] C. Carlet, C. Ding, and J. Yuan, "Linear codes from perfect nonlinear mappings and their secret sharing schemes," *IEEE Transactions on Information Theory*, vol. 51, pp. 2089–2102, 2005. [↗](#)
- [9] Z. Zhou and C. Ding, "A class of three-weight cyclic codes," *Finite Fields and Their Applications*, vol. 25, pp. 79–93, 2014. [↗](#)
- [10] S. Chang and J. Hyun, "Linear codes from simplicial complexes," *Des. Codes Cryptogr*, vol. 86, pp. 2167–2181, 2018. [↗](#)
- [11] C. Ding, Z. Heng, and Z. Zhou, "Minimal binary linear codes," *IEEE Transactions on Information Theory*, vol. 64, pp. 6536–6545, 2018. [↗](#)

Public Sentiment Analysis and Distribution Optimization MBG

Akhmad Sultoni^{1*}, Dimaz Andhika Putra², Hidayah Noor
Wahidah³, Muhammad Mahathir Arief⁴, and Pelangi Jingga
Putria R.M.⁵

^{1,2,3,4,5}Department of Mathematics, University of Gadjah Mada, Yogyakarta,
Indonesia

¹akhmad.s@mail.ugm.ac.id, ²dimazandhikaputra@mail.ugm.ac.id,

³hidayah.noor.wahidah@mail.ugm.ac.id, ⁴muhammadmahathirarief@mail.ugm.ac.id,

⁵pelangi.jinggaputriarisnamaulana@mail.ugm.ac.id

Abstract. This study analyzes public sentiment and regional prioritization regarding Indonesia's Makan Bergizi Gratis (MBG) program, a national initiative aimed at reducing stunting through the distribution of free meals to schoolchildren and pregnant women. Sentiment analysis was conducted on 47,803 posts from the social media platform X (formerly Twitter) using a lexicon-based labeling method and TF-IDF feature extraction. The results show that 22,504 posts (47.1%) expressed positive sentiment, 20,010 (41.9%) negative, and 5,289 (11.0%) neutral, indicating strong public support accompanied by considerable concerns. Eleven classification models were evaluated, with the Linear Support Vector Machine (SVM) achieving the highest accuracy (96.33%), and BERT-based models also demonstrating strong performance. Latent Dirichlet Allocation (LDA) topic modeling revealed five major themes in the negative sentiment, including transparency issues, maternal and child health, and inequality of access. Furthermore, provincial-level clustering using the K-Means algorithm grouped regions into three priority levels based on socio-economic and health indicators. These findings provide critical insights for optimizing policy targeting and efficient resource allocation in the implementation of the MBG program.

Keywords: Sentiment Analysis, Makan Bergizi Gratis, Clustering, Stunting.

*Corresponding author

2020 Mathematics Subject Classification: 91C20

Received: 05-07-2025, accepted: 18-09-2025.

1. INTRODUCTION

Stunting is one of the chronic nutritional problems that has become a major concern for the Indonesian government. Based on the 2023 Indonesian Nutrition Status Survey, the government has targeted a reduction in the stunting rate to 14% by 2024, considering that its prevalence was still at 21.6% in 2022 [1]. In addition, malnutrition remains a serious public health issue in Indonesia, particularly among children. According to data from the Ministry of Health, approximately 3.8% of children under five in Indonesia experience malnutrition. This highlights a significant challenge in efforts to improve child nutrition and health. Various initiatives have been carried out to address this issue, one of which is the *Makan Bergizi Gratis* (MBG) program, formerly known as the Free Lunch Program. This program is an initiative of the President and Vice President elected in 2024, aimed at improving the nutritional intake of children and pregnant women by providing free lunches and milk at schools, pesantren (Islamic boarding schools), and for pregnant women, as part of the national stunting alleviation strategy [2].

Although the *Makan Bergizi Gratis* program offers potential benefits in efforts to combat stunting, it has also sparked controversy among the public, particularly on social media platforms such as X (formerly Twitter). One of the main issues under scrutiny is the substantial budget required, which is estimated to reach Rp450 trillion. Furthermore, the proposed funding plan—reportedly involving the use of the State Budget (APBN) for the education sector and School Operational Assistance (BOS) funds—has triggered opposition, including from the Indonesian Teachers Union Federation (FSGI) [3]. Concerns have also been raised regarding the potential impact of this policy on education costs, teacher welfare, and the overall financial stability of the country. Amid this public debate, doubts have emerged about the program’s feasibility and long-term sustainability [4]. The growing polarization of public opinion highlights the need for sentiment analysis to better understand public perceptions of the policy, which could serve as valuable input for policymakers in evaluating and refining the program.

Sentiment analysis is a computational process aimed at evaluating individuals’ opinions, feelings, and emotions toward an entity, event, or related attribute [5]. Its primary focus is to identify the polarity of a text—whether it is positive, negative, or neutral [6]. In today’s digital era, people are increasingly active in using social media as a medium to express their views and emotions on various issues, including public policies. Therefore, it is important to analyze public sentiment and its changes over time to gain deeper insights into the public’s responses to current issues [7]. In this study, data were collected from the social media platform X, which is widely used by users to discuss and respond to trending topics. In addition to serving as a space for expression, this platform is also considered a representative and reliable source of data for capturing public perception and opinion on ongoing policies and social phenomena, including *Makan Bergizi Gratis* [8].

In conducting sentiment classification, this study employs the Support Vector Machine (SVM) algorithm, which is known as one of the most effective and reliable

methods for text classification tasks [9]. SVM works by separating data into two or more classes through the identification of an optimal hyperplane that maximizes the margin between data points from each class [10]. In other words, the algorithm seeks the best decision boundary that minimizes classification errors. Its ability to handle high-dimensional data and its robustness against overfitting make SVM particularly suitable for text-based sentiment analysis, especially when dealing with social media data, which tends to be diverse and unstructured.

In addition to analyzing public sentiment, this study also aims to identify the level of need for the *Makan Bergizi Gratis* (MBG) program across various provinces in Indonesia. The identification process utilizes data from the Central Bureau of Statistics (BPS), which includes key indicators such as population size, stunting prevalence, poverty rate, average income, and education level. To classify the provinces based on their level of need, a clustering method is employed, allowing the division of regions into three categories: high, moderate, and low need. The results of this clustering are expected to provide a solid foundation for determining priority target areas, enabling the design of policies that are more focused, well-targeted, and efficient in accelerating the national effort to reduce stunting.

2. Related Work

Previous studies on sentiment analysis of users on platform X toward the *Makan Bergizi Gratis* (MBG) program include works by [11], [12], and [8]. These studies were limited to the use of three classification algorithms: Support Vector Machine (SVM), Random Forest, and Naive Bayes Classifier. In contrast, our study not only employs a broader range of machine learning models for sentiment classification, but also introduces a regional clustering analysis to identify variations in MBG-related needs across different provinces in Indonesia.

3. Methodology

3.1. Data Collection. This study utilizes two main datasets corresponding to the sentiment analysis and the clustering of MBG needs across provinces.

For the sentiment analysis, textual data were collected from the social media platform X (formerly Twitter). The data acquisition was conducted using keyword-based scraping with search terms such as “*Makan Bergizi Gratis*,” “*Program Gizi*,” “*Stunting*,” and related hashtags. The collected data consist of user-generated posts reflecting public opinion on the MBG program. The complete dataset can be accessed via the following link:

https://drive.google.com/drive/folders/10Ab9G2avR0fv_BL82uLIxPeeX0VWG6qV

For the MBG clustering analysis, we obtained structured quantitative data from the official portal of **Badan Pusat Statistik (BPS) Republik Indonesia**. The dataset includes a comprehensive set of socio-economic and health indicators at the provincial level, such as:

- Human Development Index (HDI)
- Gini Ratio
- Total Population
- Special Index for Stunting Management
- Stunting Prevalence among Children Under Five
- Open Unemployment Rate
- Number of Families at Risk of Stunting
- Percentage of Population Living Below the Poverty Line
- Poverty Depth Index
- Poverty Severity Index
- Average Hourly Wage
- Expenditure per Capita
- Prevalence of Inadequate Food Consumption

All indicators were compiled and normalized for clustering purposes. The processed dataset is available at:

https://docs.google.com/spreadsheets/d/1nFo6YTkCv-it7_EHkw2T9_BkBnBzwMpJh_rysChgEGY

3.2. Sentiment Analysis. Sentiment analysis is the process of analyzing textual data to determine its polarity, i.e., whether the opinion is positive, negative, or neutral. In this study, sentiment analysis is applied to X posts related to the *Makan Bergizi Gratis* program.

3.2.1. Preprocessing. The preprocessing steps are critical to ensure the quality of input data [5].

- **Cleaning:** Removing emojis, URLs, symbols, and punctuations.
- **Case folding:** Converting all characters to lowercase.
- **Normalization:** Replacing informal or slang words with their formal equivalents.
- **Tokenization:** Splitting text into individual words (tokens).
- **Stopword removal:** Eliminating common words that carry little semantic value.
- **Stemming:** Reducing words to their root forms using the Sastrawi stemmer.

3.2.2. Labeling. Sentiment labeling is conducted using the lexicon-based method with the InSet lexicon [13], which provides lists of positive and negative words. Each text is assigned a score:

- *Positive* if score > 0
- *Negative* if score < 0
- *Neutral* if score $= 0$

The score is calculated as:

$$\text{Sentiment Score} = \text{Positive Words} - \text{Negative Words} \quad (1)$$

3.2.3. *Feature Extraction.* Each document is transformed into a numerical vector using Term Frequency–Inverse Document Frequency (TF-IDF) [14]:

$$\text{TF-IDF}(t, d) = tf(t, d) \times \log \left(\frac{N}{df(t)} \right) \quad (2)$$

where $tf(t, d)$ represents the frequency of term t in document d ; $df(t)$ denotes the number of documents in the corpus that contain the term t ; and N is the total number of documents in the corpus. These components are used to compute the TF-IDF weight, which reflects how important a word is to a document in a given collection.

3.2.4. *Classification Modeling.* Eleven machine learning models are used:

- **Logistic Regression**

Logistic Regression is a statistical method used for binary classification problems. It models the probability that a given input x belongs to a certain class $y = 1$ using the sigmoid (logistic) function:

$$P(y = 1|x) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n)}} \quad (3)$$

where x_i denotes the i -th feature, β_i represents the coefficient for feature x_i , and β_0 is the intercept term. This method is interpretable and effective for linearly separable data [15].

- **Multinomial Naive Bayes**

Multinomial Naive Bayes is a generative model based on Bayes' theorem, widely used for text classification. It assumes that features (typically word frequencies) are conditionally independent given the class. The classification rule is given by:

$$P(c|x) = \frac{P(c) \prod_{i=1}^n P(x_i|c)}{P(x)} \quad (4)$$

where c is the class label, $x = (x_1, x_2, \dots, x_n)$ is the feature vector representing word counts or frequencies, and $P(x_i|c)$ is the likelihood of word x_i given class c . It is simple, fast, and suitable for high-dimensional problems [16, 17, 18].

- **Support Vector Machine (SVM)**

Support Vector Machine (SVM) is a powerful supervised learning model used for classification and regression tasks. It identifies the optimal hyperplane that maximally separates different class labels. For linear classification:

$$f(x) = \text{sign}(w \cdot x + b) \quad (5)$$

where w is the weight vector, x is the input feature vector, and b is the bias term. For non-linear data, kernel functions such as the radial basis function (RBF) are employed:

$$K(x_i, x_j) = \exp(-\gamma \|x_i - x_j\|^2) \quad (6)$$

with γ controlling the influence of a single training example. SVM excels in both linearly and non-linearly separable data [19, 20, 21].

- **Random Forest**

Random Forest is an ensemble learning method that constructs multiple decision trees and merges their outputs through majority voting for classification tasks:

$$\hat{y} = \text{majority_vote}(h_1(x), h_2(x), \dots, h_n(x)) \quad (7)$$

where $h_i(x)$ is the prediction of the i -th decision tree for input x . This technique improves predictive accuracy and reduces overfitting [22, 23, 24].

- **AdaBoost**

AdaBoost, or Adaptive Boosting, is an ensemble method that combines multiple weak learners in a sequential manner. Each subsequent model focuses on instances that were misclassified by previous ones. The weight of each learner is computed as:

$$\alpha_t = \frac{1}{2} \ln \left(\frac{1 - \epsilon_t}{\epsilon_t} \right) \quad (8)$$

where ϵ_t is the error rate of the t -th weak classifier, and α_t is its weight. Though sensitive to noisy data, AdaBoost can enhance model accuracy significantly [25].

- **XGBoost**

XGBoost (Extreme Gradient Boosting) is an advanced implementation of gradient boosting algorithms. It incorporates regularization and second-order derivatives for enhanced performance. The loss function at iteration t is approximated by:

$$\mathcal{L}^{(t)} \approx \sum_{i=1}^n \left[g_i f_t(x_i) + \frac{1}{2} h_i f_t(x_i)^2 \right] + \Omega(f_t) \quad (9)$$

where g_i and h_i are the first and second-order gradients of the loss with respect to predictions $f_t(x_i)$, and Ω is a regularization term. XGBoost is efficient, scalable, and robust against overfitting [26].

- **LightGBM**

LightGBM is a fast and scalable gradient boosting framework that uses histogram-based algorithms and grows trees leaf-wise with depth constraints. It is optimized for memory usage and training time, making it suitable for large-scale data processing [27].

- **BERT (Bidirectional Encoder Representations from Transformers)**

BERT is a transformer-based language model that captures the full context of a word by looking at both its left and right surroundings. The core mechanism is self-attention:

$$\text{Attention}(Q, K, V) = \text{softmax} \left(\frac{QK^T}{\sqrt{d_k}} \right) V \quad (10)$$

where Q (queries), K (keys), and V (values) are matrices derived from the input, and d_k is the dimensionality of the keys. BERT excels in a variety of NLP tasks including sentiment analysis, question answering, and more [28].

- **DistilBERT**

DistilBERT is a distilled version of BERT that retains most of its performance while being smaller and faster. It is trained using knowledge distillation to match the behavior of BERT with fewer parameters and computations, making it ideal for real-time or low-resource applications [29].

- **IndoBERTweet**

IndoBERTweet is a variant of BERT pre-trained on Indonesian tweets. It is tailored to understand informal language, slang, abbreviations, and other characteristics unique to Indonesian social media. This makes it especially effective for sentiment analysis and opinion mining in Indonesian contexts [30].

Classical models are trained on TF-IDF vectors using `scikit-learn` [31]. Transformer-based models utilize pre-trained embeddings from Indonesian language models such as IndoBERTweet, powered by HuggingFace Transformers [32].

3.2.5. *Evaluation Metrics.* Model performance is evaluated using:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (11)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (12)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (13)$$

$$\text{F1-score} = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (14)$$

3.3. **Topic Modeling.** Topic modeling is an unsupervised learning technique used to discover hidden thematic structures in text corpora.

3.3.1. *Latent Dirichlet Allocation (LDA).* Topic analysis was conducted using the Latent Dirichlet Allocation (LDA) method, an unsupervised learning algorithm commonly used to uncover latent thematic structures within a collection of documents [33]. LDA operates under the assumption that each document is a mixture of multiple topics, and each topic is characterized by a specific distribution over words.

Prior to model training, data categorized as negative sentiment underwent a series of preprocessing steps, including case folding, removal of non-alphabetic characters, tokenization, and stopword removal using an Indonesian stopword dictionary. The resulting tokens were then converted into a bag-of-words representation and further processed into a corpus and dictionary using the `Gensim` library [34].

The probability of a topic $z_i = k$ for word w_i in document d_i is given by:

$$P(z_i = k | w_i = w, d_i = d) \propto \frac{n_{dk}^{-i} + \alpha}{n_d^{-i} + K\alpha} \cdot \frac{n_{kw}^{-i} + \beta}{n_k^{-i} + V\beta} \quad (15)$$

where n_{dk}^{-i} denotes the number of words in document d that are assigned to topic k , excluding the current word i ; n_{kw}^{-i} represents the number of times word w is assigned to topic k , also excluding the current word. The parameters α and β are Dirichlet hyperparameters that control the sparsity of topic and word distributions, respectively. K refers to the total number of latent topics assumed in the model, while V denotes the vocabulary size, or the number of unique terms in the corpus. These parameters are used in the Gibbs sampling update equation for estimating the topic distribution in Latent Dirichlet Allocation.

The LDA model was trained using five topics and ten passes (iterations), producing outputs in the form of dominant keywords for each topic, along with topic distributions across the documents. Model visualization was carried out using the `pyLDAvis` library to enable interactive interpretation of inter-topic relationships.

3.4. Clustering Analysis.

3.4.1. *K-Means Clustering*. **Basic Concept of K-Means**

K-Means clustering is one of the most widely used unsupervised learning algorithms for data grouping [35]. Its main objective is to partition a set of n data points into k distinct clusters based on similarity. The algorithm aims to assign each observation to the cluster with the nearest mean, known as the *centroid*, thereby minimizing intra-cluster variance. The centroid is calculated as the average of all data points in a given cluster. K-Means thus seeks to minimize the total within-cluster variation, also known as the sum of squared errors (SSE), ensuring each cluster is as homogeneous as possible.

Algorithm Steps

The K-Means algorithm typically follows these steps [36, 37]:

- (1) **Initialization:** Choose the number of clusters k to form.
- (2) **Initial Centroids:** Randomly select k initial centroids from the dataset.
- (3) **Assignment Step:** Assign each data point to the nearest centroid using a distance metric, commonly Euclidean distance.
- (4) **Update Step:** Recalculate the centroid of each cluster by computing the mean of all points assigned to that cluster.
- (5) **Convergence Check:** Repeat steps 3 and 4 until convergence, i.e., when there are no further changes in cluster assignments or centroids.

Mathematical Formulation

Euclidean Distance

The Euclidean distance between a data point $\mathbf{x}$ and a centroid $\boldsymbol{\mu}$ in d -dimensional

space is calculated as:

$$d(\mathbf{x}, \boldsymbol{\mu}) = \sqrt{\sum_{j=1}^d (x_j - \mu_j)^2} \quad (16)$$

Cluster Centroid Calculation

The centroid μ_k of cluster k is computed as the mean of all N_k points in that cluster [37]:

$$\mu_k = \frac{1}{N_k} \sum_{x \in S_k} x \quad (17)$$

where S_k is the set of data points in cluster k , and $N_k = |S_k|$ is the number of points in that cluster.

Objective Function (SSE)

The objective of K-Means is to minimize the total sum of squared distances between data points and their respective cluster centroids [36]:

$$J = \sum_{k=1}^K \sum_{x \in S_k} \|x - \mu_k\|^2 \quad (18)$$

This function, known as the within-cluster sum of squares (WCSS), quantifies the compactness of the clusters. The algorithm stops when J converges to a local minimum.

3.4.2. K-Median Clustering. Basic Concept

K-Median clustering is an alternative to K-Means that uses median values instead of means to define cluster centers. It is particularly useful when the dataset contains outliers or is not normally distributed. Unlike K-Means, which minimizes the sum of squared Euclidean distances, K-Median minimizes the sum of absolute distances (Manhattan distances) between data points and the cluster centroids (medians). This makes K-Median more robust to noise and extreme values [38].

Algorithm Steps

The K-Median algorithm proceeds as follows:

- (1) **Initialize:** Choose the number of clusters k and randomly select k initial medians.
- (2) **Assignment:** Assign each data point to the nearest median based on Manhattan distance.
- (3) **Update:** For each cluster, update the median by computing the component-wise median of the assigned points.
- (4) **Repeat:** Iterate steps 2 and 3 until cluster assignments no longer change or a convergence criterion is met.

Mathematical Formulation

Given a dataset $X = \{x_1, x_2, \dots, x_n\}$ and a set of cluster centers $\{m_1, m_2, \dots, m_k\}$,

the K-Median objective function is:

$$J = \sum_{i=1}^n \min_{j \in \{1, \dots, k\}} \|x_i - m_j\|_1 \quad (19)$$

Here, $\|\cdot\|_1$ denotes the Manhattan (L1) norm. Each cluster center m_j is updated as the median of all data points assigned to cluster j :

$$m_j = \text{median}\{x_i \mid x_i \in C_j\} \quad (20)$$

Comparison with K-Means

While K-Means minimizes the sum of squared Euclidean distances and is sensitive to outliers, K-Median is more robust by minimizing the sum of absolute deviations. This makes K-Median particularly effective for applications involving skewed or noisy data.

3.4.3. Clustering Analysis with Genetic Algorithm. Genetic algorithm is a search and optimization method inspired by the principles of natural selection and biological genetics. This approach begins by generating a number of random solutions that form a population of chromosomes. Through evolutionary stages—including selection, crossover, and mutation—the algorithm aims to find a globally optimal solution. Selection is carried out by choosing chromosomes with the highest fitness values to form a new generation. The crossover process then combines genetic information from two parents to produce offspring with superior characteristics [39]. Genetic algorithm is applied to K-Means clustering to enhance clustering performance. The performance of K-Means clustering is known to be sensitive to suboptimal initial centroid selection. By using a genetic algorithm, the search for more representative cluster centers can be conducted more thoroughly.

In the implementation of genetic algorithm for K-Means clustering, the process begins by forming an initial population consisting of candidate solutions in the form of different centroid positions. Each chromosome represents a set of centroids as a potential solution. Evaluation of each chromosome is done by calculating its fitness value, typically using the within-cluster sum of squares, which measures how well the centroids divide the data. Selection then chooses the best-performing chromosomes to generate the next generation, followed by a crossover process that combines characteristics from two parent solutions to produce improved offspring. Random mutation is applied to maintain population diversity and prevent convergence to local optima. The processes of selection, crossover, and mutation are repeated over several generations until convergence is achieved or the best fitness solution is found, resulting in more optimal initial centroids for clustering.

3.5. Experimental Setup. This study consists of two main experimental components: sentiment analysis and provincial clustering.

3.5.1. Sentiment Analysis. To capture public response toward the Makan Bergizi Gratis (MBG) program, we collected textual data from social media platform **X**

(formerly Twitter) using relevant keywords and hashtags (e.g., #MBG, #Makan-BergiziGratis, #ProgramGizi). The dataset was preprocessed using standard NLP techniques such as case-folding, tokenization, stopword removal, and stemming.

We implemented a supervised machine learning model for sentiment classification, categorizing each post as *positive*, *negative*, or *neutral*. We used the Python `scikit-learn` library, employing a TF-IDF vectorizer for feature extraction and a Logistic Regression classifier. Hyperparameters were tuned using grid search with 5-fold cross-validation. Performance metrics such as accuracy, precision, recall, and F1-score were evaluated on a 20% hold-out test set.

3.5.2. Clustering of Provincial MBG Needs. In the second phase, we analyzed official provincial-level indicators from Badan Pusat Statistik (BPS), including stunting prevalence, poverty rate, population, average income, and education level. The data were normalized using Min-Max scaling.

We applied K-Means clustering to group provinces based on their relative need for the MBG program. The optimal number of clusters (k) was determined using the Elbow Method and Silhouette Coefficient. Provinces were then categorized into three priority levels: *high need*, *moderate need*, and *low need*.

3.5.3. Computational Environment. All experiments were conducted using Python on the **Google Colab** platform, which provides cloud-based computation and interactive visualization capabilities. The main libraries used include `pandas` for data manipulation, `scikit-learn` for machine learning modeling, and `matplotlib` and `seaborn` for data visualization.

Sentiment data was collected from the social media platform **X** (formerly Twitter) via its API using keywords and hashtags related to the MBG program. Meanwhile, provincial-level indicator data was obtained from the official website of Badan Pusat Statistik (BPS) and manually compiled into a structured dataset using Google Sheets. The dataset can be accessed at the following link:

https://docs.google.com/spreadsheets/d/1nFo6YTkCv-it7_EHkw2T9_BkBnBzwMpJh_rysChgEGY/edit?usp=sharing

All source code and experimental documentation are publicly available for reproducibility at the following links:

- **Sentiment Analysis:** https://drive.google.com/drive/folders/10Ab9G2avR0fv_BL82uLIxPeeX0VWG6qV?usp=drive_link
- **MBG Clustering by Province:** <https://colab.research.google.com/drive/1EJgCbXpF8VIppbvfcMVghH3ZrGQfYC6A?usp=sharing>

4. Results and Discussion

4.1. Sentiment Analysis.

4.1.1. Data Preprocessing. The preprocessing stage begins with data cleaning, which involves the removal of emojis, symbols, URLs, and irrelevant punctuation. This is followed by case folding to standardize all characters to lowercase and the normalization of informal or slang words into their formal equivalents. The next step is

tokenization, where sentences are segmented into individual words. Subsequently, stopwords removal is applied to eliminate common words that carry little semantic weight in the context of the analysis. Finally, stemming is performed using the Sastrawi library to reduce words to their root forms, thereby enhancing the quality of textual features used in the classification model.

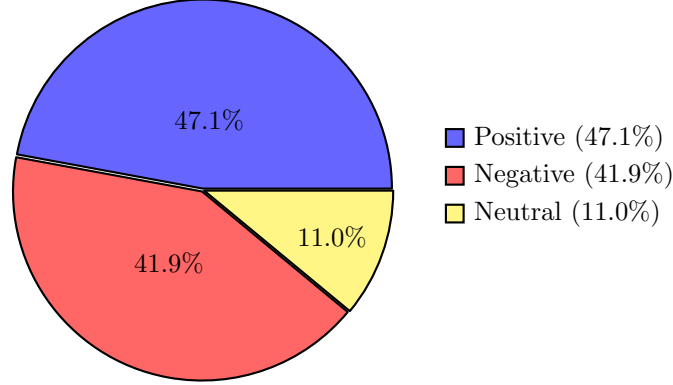


FIGURE 1. Sentiment distribution in the dataset. The legend shows each sentiment label and its corresponding percentage.

4.1.2. *Labeling.* Sentiment labeling was conducted using a lexicon-based approach, categorizing each text into positive, negative, or neutral sentiment. As shown in Figure (1), the dataset contains 22,504 positive samples (47.1%), 20,010 negative samples (41.9%), and 5,289 neutral samples (11.0%). This distribution indicates a generally positive trend in public sentiment, although a substantial portion also expresses negative opinions, suggesting ongoing debates or concerns regarding the topic.

4.1.3. *Feature Extraction.* Feature extraction was performed using the Term Frequency Inverse Document Frequency (TF-IDF) method, with a maximum of 5,000 features. This method generates a numerical representation of the text by quantifying the importance of each word in a document relative to the entire corpus.

4.1.4. *Modeling.* The dataset was divided into training and testing sets in an 80:20 ratio using stratified sampling to preserve the distribution of sentiment classes. A total of eleven classification models were employed in this study, comprising eight traditional machine learning models—Logistic Regression, Multinomial Naive Bayes, Support Vector Machines (with both Linear and RBF kernels), Random Forest, AdaBoost, XGBoost, and LightGBM—and four transformer-based models: BERT, DistilBERT, BERTweet, and IndoBERTweet.

4.1.5. *Model Evaluation.* Performance comparison among 11 sentiment classification models is presented in Table (I), covering both traditional machine learning and transformer-based approaches. The metrics evaluated include accuracy, precision, recall, and F1-score for both positive and negative classes.

TABLE 1. Performance Evaluation of Sentiment Classification Models

Model	Acc	P _{Pos}	R _{Pos}	F1 _{Pos}	P _{Neg}	R _{Neg}	F1 _{Neg}
Logistic Regression	0.9410	0.9496	0.9382	0.9439	0.9315	0.9440	0.9377
Multinomial NB	0.7976	0.8018	0.8205	0.8110	0.7927	0.7719	0.7821
SVM (Linear)	0.9633	0.9653	0.9653	0.9653	0.9610	0.9610	0.9610
SVM (RBF)	0.9467	0.9552	0.9436	0.9494	0.9374	0.9503	0.9438
Random Forest	0.8523	0.8603	0.8607	0.8605	0.8433	0.8428	0.8430
AdaBoost	0.6953	0.7870	0.5819	0.6691	0.6363	0.8228	0.7177
XGBoost	0.8578	0.8893	0.8354	0.8615	0.8267	0.8831	0.8539
LightGBM	0.8745	0.8978	0.8609	0.8790	0.8505	0.8898	0.8697
BERT	0.9153	0.9219	0.9178	0.9198	0.9080	0.9125	0.9103
DistilBERT	0.9066	0.9125	0.9109	0.9117	0.9000	0.9018	0.9009
BERTweet	0.9013	0.9148	0.8971	0.9059	0.8868	0.9060	0.8963
IndoBERTweet	0.8752	0.8914	0.8703	0.8807	0.8579	0.8808	0.8692

Among all models, Support Vector Machine with linear kernel (SVM Linear) achieved the highest overall performance, with an accuracy of 96.33%, and balanced precision, recall, and F1-score for both sentiment classes (all exceeding 96%). This indicates strong generalization and consistency in detecting sentiment polarity from social media text.

Transformer-based models also performed competitively. BERT achieved an accuracy of 91.53%, closely followed by DistilBERT and BERTweet, with F1-scores above 90% for both classes. These results affirm the effectiveness of pre-trained language models in capturing nuanced sentiment in informal and context-rich data.

In contrast, traditional models such as Multinomial Naive Bayes and AdaBoost demonstrated lower accuracy, at 79.76% and 69.53% respectively, with significantly imbalanced performance between positive and negative classes. This highlights their limitations in handling the complexity of social media text, particularly with regard to sarcasm, slang, and varying sentence structures.

Overall, the evaluation suggests that while classical models like SVM Linear remain highly effective with well-engineered features (e.g., TF-IDF), transformer-based models offer robust alternatives for future work, particularly when dealing with larger and more diverse datasets.

4.2. Topic Modelling.

4.2.1. *LDA Topic Modeling.* The results of LDA analysis indicate that negative sentiment toward the free meal program can be grouped into five major themes:

(1) **Children’s Education and Nutrition**

Criticisms highlight disparities in access to child nutrition programs, particularly in remote regions such as Papua. Dominant keywords: *susu*, *anak*, *bantu*, *Papua*.

(2) **Maternal and Infant Welfare**

Complaints focus on the lack of government support for pregnant women

and infants, often accompanied by sarcastic remarks directed at public officials. Dominant keywords: *hamil, balita, sejahtera, hidup*.

(3) **Program and Funding Transparency**

Criticisms address the lack of clarity regarding the distribution of information and funds. Dominant keywords: *uang, hilang, informasi, masak*.

(4) **Budgeting and Public Policy Implementation**

Issues related to national budget (APBN) management, meal quality, and program implementation in schools. Dominant keywords: *menu, orang tua, budget, apbn*.

(5) **Public Distrust of Government**

Negative and often harsh expressions toward government programs, reflecting widespread public distrust. Dominant keywords: *tolol, tanggung, duit, gratis*.

4.3. Clustering Analysis. In this study, clustering analysis is applied to data related to health and nutrition, social and demographic conditions, as well as economic and employment indicators for each province in Indonesia. The data used is from 2023 and sourced from Badan Pusat Statistik Indonesia. The analysis focuses on data from 34 provinces in Indonesia, excluding the four new provinces established in 2022, as data for several variables in these provinces is not yet available. The purpose of the clustering analysis is to categorize regions based on their level of priority for receiving Makan Bergizi Gratis (MBG) program. This will enable the government to focus the implementation of the program on high-priority areas, ensuring more effective and targeted resource allocation. Priority levels are determined based on demographic factors, economic conditions, and the health and nutritional status of each region, with the goal of optimizing government spending for the program's implementation.

The clustering analysis is conducted using several algorithms, including K-Means, K-Median, and K-Means with Genetic Algorithm Optimization. The analysis aims to form three clusters, as predetermined by the researchers, representing high, medium, and low priority groups. The results from the three algorithms are then compared using the silhouette score evaluation metric to identify the most effective algorithm for clustering provinces based on their priority level. The clustering outcomes for each algorithm are presented in the following Table (2).

TABLE 2. Clustering Performance

Algorithm	Silhouette Score
K-Means	0.6047
K-Median	0.4266
K-Means with GA	0.5729

Based on the clustering results from the three algorithms, the K-Means algorithm demonstrated the best performance, achieving the highest silhouette score compared to K-Median and K-Means with Genetic Algorithm Optimization. Therefore, the clustering results from K-Means will be used to determine the priority

levels of each region. Table (3) and Figure (2).presents the provincial clustering outcomes using the K-Means algorithm.

TABLE 3. Clustering Analysis Results

Cluster	Province
Cluster 1	Aceh, Sumatera Utara, Sumatera Barat, Jambi, Sumatera Selatan, Bengkulu, Lampung, Jawa Tengah, Jawa Timur, NTT, NTB, Kalimantan Barat, Kalimantan Selatan, Sulawesi Utara, Sulawesi Tengah, Sulawesi Selatan, Sulawesi Tenggara, Gorontalo, Sulawesi Barat, Maluku, Maluku Utara
Cluster 2	Riau, Kep. Bangka Belitung, Kep. Riau, Jawa Barat, DI Yogyakarta, Banten, Bali, Kalimantan Tengah, Kalimantan Timur, Kalimantan Utara, Papua Barat, Papua
Cluster 3	DKI Jakarta

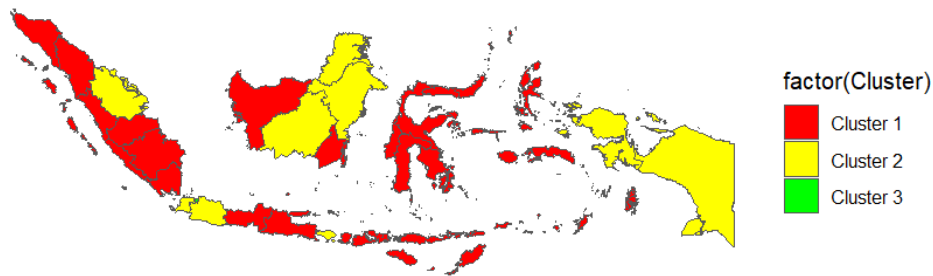


FIGURE 2. Provincial clustering results of the MBG program in Indonesia: Cluster 1 (red), Cluster 2 (yellow), and Cluster 3 (green).

Next, to identify the priority level of each cluster, it will be determined based on the characteristics of each cluster as observed from the centroid values of each variable within each cluster. The following are the centroids for each cluster in Table (4).

Provinces in Cluster 1 still show a quality of life that is not yet optimal, with the Human Development Index (HDI) categorized as moderate. The socioeconomic conditions in this area face significant challenges, as evidenced by the high

prevalence of stunting and poverty that still threaten many families. Additionally, the community's purchasing power and wage levels are relatively low, while income inequality is at a moderate level. Therefore, this region requires special attention in social and public health development efforts to significantly improve the quality of life of its population.

TABLE 4. Centroids Each Cluster

Variables	Cluster 1	Cluster 2	Cluster 3
IPM	71.672	74.270	82.460
GR	0.338	0.349	0.431
Populasi	8301.162	7808.300	10672.100
IKPS	70.624	67.883	73.900
PBS	24.100	19.925	17.600
TPT	4.378	4.867	6.530
KRS	207.098	33.146	32.850
PPM	10.878	9.180	4.440
IKdM	1.804	1.695	0.690
IKpM	0.445	0.493	0.170
RRU	17127.905	22205.750	42354.000
PpK	1257437.670	1694760.080	2791716.000
PKK	11.247	12.654	2.570

Provinces in Cluster 2 show better development compared to Cluster 1, with HDI and quality of life relatively improved. Poverty levels and stunting prevalence in this area have started to decline, resulting in a much smaller number of families at risk of stunting. Purchasing power and wages have improved, although there are still challenges related to insufficient consumption among some groups. Then, cluster 3 represents regions that have reached a very advanced level of development, with a very high HDI reflecting excellent quality of education, health, and living standards. In this area, poverty and social risks are very low, while purchasing power is relatively high. The prevalence of stunting is also very low, indicating the success of various health and social programs implemented. Nevertheless, these regions still face challenges such as unemployment and income inequality that need to be managed well.

Based on the centroids of each cluster representing the economic conditions as well as the health and nutritional status of each region, priority levels for the implementation of the Makan Bergizi Gratis (MBG) program can be determined. Cluster 1 consists of provinces with a high priority for receiving Makan Bergizi Gratis (MBG) program. This is because these areas require special attention in social and public health development efforts to significantly improve the quality of life of their populations. Cluster 2 falls under medium priority for Makan Bergizi Gratis (MBG) program, where the government can target vulnerable groups such as poor families or children in certain schools. Meanwhile, Cluster 3 includes provinces with a low priority for Makan Bergizi Gratis (MBG) program, such as DKI Jakarta.

In this cluster, implementing Makan Bergizi Gratis (MBG) program is not yet an urgent priority; the government can focus more on lower-budget initiatives such as nutrition education. By referring to the results of this clustering analysis, the government can reassess the implementation of Makan Bergizi Gratis (MBG) program, allowing it to be carried out more targetedly and to save budget.

4.4. Limitations. This study has several limitations that should be acknowledged. First, the sentiment analysis was conducted solely on the social media platform **X** (formerly Twitter). While this platform provides timely and high-volume user-generated content, it does not capture sentiments from other widely used platforms such as Facebook, Instagram, or TikTok, which may reflect different user demographics and engagement patterns. As such, the sentiment findings may not be fully representative of the broader public opinion regarding the MBG program.

Second, the clustering analysis of MBG needs was performed at the provincial level due to data availability and granularity. While this provides a general overview of regional disparities, it lacks the precis

5. CONCLUSION

Sentiment analysis reveals that the majority of responses to the MBG program are positive (47.1%), followed by negative (41.9%) and neutral (11.0%), indicating strong public support, albeit with notable concerns. Among the 11 classification models evaluated, Linear SVM achieved the highest accuracy (96.33%) with balanced performance. Transformer-based models such as BERT and DistilBERT also performed well, effectively capturing the nuances of social media language. In contrast, traditional models like Naive Bayes and AdaBoost showed lower accuracy and less consistent performance across sentiment classes. Overall, transformer-based models are a strong choice for future analysis, particularly when dealing with complex and informal social media data.

The clustering analysis classified 34 provinces into three priority levels for the Makan Bergizi Gratis (MBG) program, with K-Means showing the best performance. Cluster 1 (high priority) includes provinces with lower development indicators, while Cluster 2 and Cluster 3 represent medium and low priority, respectively. These results enable more targeted and efficient program implementation. Future research is encouraged to use more granular data at the district or sub-district (kecamatan) level to improve policy targeting and resource allocation.

REFERENCES

- [1] K. K. RI, "Prevalensi stunting di indonesia turun ke 21,6% dari 24,4%," *Journal Name*, 2024.

- [2] T. . Rachmawati, "Implementasi algoritma naive bayes untuk analisis sentimen terhadap program makan siang gratis," *Jurnal Indonesia: Manajemen Informatika dan Komunikasi*, vol. 5, no. 3, pp. 2925–2939, 2024.
- [3] N. A. Rahmawati, S. A. Prasetyo, and M. W. Ramadhani, "Memetakan visi prabowo gibran pada masa kampanye dalam prespektif pembangunan:(analisis wacana kritis visi dan misi prabowo gibran dalam prespektif moderenisasi)," *WISSEN: Jurnal Ilmu Sosial dan Humaniora*, vol. 2, no. 3, pp. 97–120, 2024.
- [4] P. A. Maharani, A. R. Namira, and T. V. Chairunnisa, "Peran makan siang gratis dalam janji kampanye prabowo gibran dan realisasinya," *Journal Of Law And Social Society*, vol. 1, no. 1, pp. 1–10, 2024.
- [5] C. Steven and W. Wella, "The right sentiment analysis method of indonesian tourism in social media twitter," *IJNMT (International Journal of New Media Technology)*, vol. 7, no. 2, pp. 102–110, 2020.
- [6] F. R. B. Kahi *et al.*, "Analisis sentimen masyarakat di twitter terhadap pemerintahan anies baswedan menggunakan metode naive bayes classifier," *Jurnal Minfo Polgan*, vol. 13, no. 1, pp. 324–336, 2024.
- [7] M. Iqbal, M. Afdal, and R. Novita, "Implementasi algoritma support vector machine untuk analisa sentimen data ulasan aplikasi pinjaman online di google play store: Implementation of support vector machine algorithm for sentiment analysis of online loan application review data on google play store," *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 4, pp. 1244–1252, 2024.
- [8] A. Sitanggang, Y. Umaidah, and R. I. Adam, "Analisis sentimen masyarakat terhadap program makan siang gratis pada media sosial x menggunakan algoritma naive bayes," *Jurnal Informatika Dan Teknik Elektro Terapan*, vol. 12, no. 3, 2024.
- [9] S. D. Wahyuni and R. H. Kusumodestoni, "Optimalisasi algoritma support vector machine (svm) dalam klasifikasi kejadian data stunting," *Bulletin of Information Technology (BIT)*, vol. 5, no. 2, pp. 56–64, 2024.
- [10] A. M. Siregar, "Analisis sentimen pindah ibu kota negara (ikn) baru pada twitter menggunakan algoritma naive bayes dan support vector machine (svm)," *Faktor Exacta*, vol. 16, no. 3, 2023.
- [11] E. Triningsih, "Analisis sentimen terhadap program makan bergizi gratis menggunakan algoritma machine learning pada sosial media x," *BUILDING OF INFORMATICS, TECHNOLOGY AND SCIENCE (BITS)*, vol. 6, no. 4, 2025.
- [12] A. Ramadhani, I. Permana, M. Afdal, and M. Fronita, "Analisis sentimen tanggapan publik di twitter terkait program kerja makan siang gratis prabowo-gibran menggunakan algoritma naive bayes classifier dan support vector machine," *Building of Informatics, Technoogy and Science (BITS)*, vol. 6, no. 3, pp. 1509–1516, 2024.
- [13] R. F. Fajri, M. Adriani, *et al.*, "Inset: A sentiment lexicon for indonesian social media," in *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, 2018.
- [14] D. Jurafsky and J. H. Martin, *Speech and Language Processing (3rd ed. draft)*. Stanford University, 2021. Available at <https://web.stanford.edu/~jurafsky/slp3/>.
- [15] D. W. Hosmer, S. Lemeshow, and R. X. Sturdivant, *Applied Logistic Regression*. Wiley, 2013.
- [16] A. McCallum and K. Nigam, "A comparison of event models for naive bayes text classification," *AAAI-98 workshop on learning for text categorization*, 1998.
- [17] A. Sabrani and F. Bimantoro, "Multinomial naive bayes untuk klasifikasi artikel online tentang gempa di indonesia," *Jurnal Teknologi Informasi, Komputer, Dan Aplikasinya (JTIKA)*, vol. 2, no. 1, pp. 89–100, 2020.
- [18] A. W. Syaputri, E. Irwandi, and Mustakim, "Naive bayes algorithm for classification of student major's specialization," *Journal not specified*, vol. 1, no. 1, 2020.
- [19] C. Cortes and V. Vapnik, "Support-vector networks," *Machine learning*, vol. 20, no. 3, pp. 273–297, 1995.

- [20] Saikin, S. Fadli, and M. Ashari, "Optimization of support vector machine method using feature selection to improve classification result," *Journal not specified*, vol. 4, no. 1, 2021.
- [21] N. H. Ovirianti, M. Zarlis, and H. Mawengkang, "Support vector machine using a classification algorithm," *Journal not specified*, vol. 7, no. 3, 2022.
- [22] L. Breiman, "Random forests," *Machine learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [23] A. E. Maxwell, T. A. Warner, and F. Fang, "Implementation of machine-learning classification in remote sensing: An applied review," *International Journal of Remote Sensing*, vol. 39, no. 9, pp. 2784–2817, 2018.
- [24] Y. Religia, A. Nugroho, and W. Hadikristanto, "Analisis perbandingan algoritma optimasi pada random forest untuk klasifikasi data bank marketing," *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 1, pp. 187–192, 2021.
- [25] Y. Freund and R. E. Schapire, "A decision-theoretic generalization of on-line learning and an application to boosting," *Journal of computer and system sciences*, vol. 55, no. 1, pp. 119–139, 1997.
- [26] T. Chen and C. Guestrin, "Xgboost: A scalable tree boosting system," *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 785–794, 2016.
- [27] G. Ke, Q. Meng, T. Finley, T. Wang, W. Chen, W. Ma, Q. Ye, and T.-Y. Liu, "Lightgbm: A highly efficient gradient boosting decision tree," in *Advances in neural information processing systems*, pp. 3146–3154, 2017.
- [28] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding," *NAACL*, 2019.
- [29] V. Sanh, L. Debut, J. Chaumond, and T. Wolf, "Distilbert, a distilled version of bert: smaller, faster, cheaper and lighter," *arXiv preprint arXiv:1910.01108*, 2019.
- [30] I. Alfina, R. Mulia, and M. I. Fanany, "Indobertweet: Pretrained transformer-based language model for indonesian twitter," *arXiv preprint arXiv:2109.05238*, 2021.
- [31] F. Pedregosa, G. Varoquaux, A. Gramfort, *et al.*, "Scikit-learn: Machine learning in python," 2011. *Journal of Machine Learning Research*, 12, 2825–2830.
- [32] T. Wolf, L. Debut, V. Sanh, *et al.*, "Transformers: State-of-the-art natural language processing." *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, 2020. HuggingFace Transformers library.
- [33] D. M. Blei, A. Y. Ng, and M. I. Jordan, "Latent dirichlet allocation," *Journal of Machine Learning Research*, vol. 3, pp. 993–1022, 2003.
- [34] R. Řehůřek and P. Sojka, "Software framework for topic modelling with large corpora," 2010. Gensim Python library.
- [35] J. MacQueen, "Some methods for classification and analysis of multivariate observations," in *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, vol. 1, pp. 281–297, University of California Press, 1967.
- [36] S. P. Lloyd, "Least squares quantization in pcm," *IEEE Transactions on Information Theory*, vol. 28, no. 2, pp. 129–137, 1982.
- [37] A. K. Jain, M. N. Murty, and P. J. Flynn, "Data clustering: A review," *ACM Computing Surveys*, vol. 31, no. 3, pp. 264–323, 1999.
- [38] V. Arya, N. Garg, R. Khandekar, A. Meyerson, K. Munagala, and V. Pandit, "Local search heuristics for k-median and facility location problems," *SIAM Journal on Computing*, vol. 33, no. 3, pp. 544–562, 2004.
- [39] M. Jahandideh-Tehrani, G. Jenkins, and F. Helfer, "A comparison of particle swarm optimization and genetic algorithm for daily rainfall-runoff modelling: a case study for southeast queensland, australia," *Optimization and Engineering*, vol. 22, no. 1, pp. 29–50, 2021.

Optimizing Linear Models with Deep Neural Networks: A Case Study on Gambling Data Across Indonesian Provinces

Astri Ayu Nastiti^{1*}, Abdurakhman²

^{1,2}Department of Mathematics, University of Gadjah Mada, Yogyakarta, Indonesia

¹astriayunastiti@mail.ugm.ac.id, ²rachmanstat@ugm.ac.id

Abstract. Linear regression analysis is a common method that are free to vary and are subject to error. In this study we used hybrid of linear regression and its family to Deep Neural Network (DNN) to fill these gaps. In this paper analyze the phenomenon of gambling in Indonesia in 2018. Results show that the hybrid model is significantly superior to the single model, with the hybrid linear model reducing RMSE by 15.9% and MAPE by 16.2% compared to the single linear model. The hybrid ridge model showed small but consistent improvements in RMSE and MAPE. The most notable improvement was seen in the hybrid lasso model which reduced RMSE by 34.1% and MAPE by 47.1% over the single lasso model. The hybrid elastic net model also showed improved performance with a decrease in RMSE by 16.9% and MAPE by 18.3%. In conclusion, the integration of traditional regression methods with DNN in this hybrid model offers a significant improvement in prediction accuracy, making it a more effective and efficient tool in the analysis of gambling phenomena.

Keywords: Linear regression, ridge regression, LASSO regression, elastic net regression, Deep Neural Network, Hybrid Model.

1. INTRODUCTION

Gambling is a social phenomenon that can have a negative impact on society, both in economic and social terms. Several studies reveal the effects of gambling including divorce [1], increased anxiety levels [2], affecting not only adults but also children [3]. Therefore, understanding the characteristics of the community and the factors that contribute to the incidence rate of gambling is an important step towards formulating effective policies to address it.

*Corresponding author

2020 Mathematics Subject Classification: 62J05, 62J07, 92B20

Received: 25-06-2025, accepted: 07-08-2025.

This paper uses population data by province in Indonesia to analyze the relationship between various socio-economic factors and the incidence of gambling. Given the relatively small amount of data, linear regression was chosen as the main analysis method. Linear regression is an appropriate choice because it is simple and effective in handling small datasets without requiring complex training and testing processes as in machine learning techniques.

In regression analysis, multiple linear regression is often used to understand the relationship between independent variables and dependent variables. However, to overcome the problem of multicollinearity and to improve the accuracy of the model, extensions of simple linear regression such as Ridge, Lasso, and Elastic Net regression have been introduced [4].

Ridge regression addresses multicollinearity by adding an L_2 penalty to the coefficients, thus preventing the coefficients from becoming too large [5]. Lasso regression introduces the L_1 penalty, which not only addresses multicollinearity but also performs feature selection by reducing some coefficients to zero, thus simplifying the model (Vidaurre et al., 2013). Elastic Net Regression is a combination of L_1 and L_2 penalties, which allows handling multicollinearity while performing feature selection, providing more flexibility in building robust and interpretable predictive models [6]. These three methods are known as an “extended linear regression family” that offers more advanced solutions for complex data analysis and potentially better predictive performance.

As is well known, linear regression and extended linear regression approaches can give results that differ from the actual data [7]. In some cases, simple linear regression models may fail to capture the complexity of the data patterns, especially when multicollinearity is present or when the relationship between variables is not strictly linear. Furthermore, to overcome the weaknesses of these traditional regression approaches, deep-linear regression hybrid artificial neural networks are used. This hybrid approach combines the analytical power of linear regression with the capability of artificial neural networks to recognize non-linear patterns and complex interactions in the data.

The purpose of this research is to explore and maximize the potential of linear regression and its families (ridge, lasso, and elastic net) in producing accurate predictions by combining them with deep learning networks. This research aims to identify the advantages of the hybrid approach in reducing prediction error compared to the use of a single model, as well as to develop predictive models that are more robust and efficient in handling data complexity. Thus, this research hopes to make a significant contribution to more accurate and reliable predictive modeling through the integration of traditional regression methods and deep learning technology.

2. LITERATURE REVIEW

2.1. Linear Regression. Linear regression is an equation model that explains the correlation of one response variable (Y) with two or more predictor variables ($X_1, X_2, \dots, X_p$). In addition, it is used to determine the direction of the relationship between the response variable and the predictor variables. The relationship between the response variable and the predictor variables is expressed as follows:

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \quad (1)$$

with:

$Y = n \times 1$ vector of dependent variables,

$X = n \times (p - 1)$ matrix of independent variables,

$\boldsymbol{\varepsilon} = n \times 1$ vector of independent normal random variables with expectation

$$E(\boldsymbol{\varepsilon}) = 0 \text{ and variance-covariance matrix } \sigma^2(\boldsymbol{\varepsilon}) = \sigma^2 I.$$

According to Firdaus (2004) [8], the least squares method or also called the Ordinary Least Square (OLS) method is one of the most popular methods in estimating linear regression models that produce the minimum number of squared errors. This method was first used by Carl Friedrich Gauss in the calculation of astronomical problems. The practical advantages of this method increased after the development of electronic computers, the formulation of calculation techniques in matrix notation, and the connection of the least squares concept to statistics.

Definition 2.1. Let $p \geq 1$ be a real number. The p -norm of vector $x = (x_1, x_2, \dots, x_n)$ is

$$\|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{(1/p)}. \quad (2)$$

For $p = 1$, we get the taxicab/manhattan norm, for $p = 2$ we get the Euclidean norm, and as p approaches ∞ the p -norm approaches the infinity norm [9].

From Equation (1) is obtained

$$\begin{aligned} \boldsymbol{\varepsilon} &= \mathbf{Y} - \mathbf{X}\boldsymbol{\beta} \\ \mathbf{S}(\boldsymbol{\beta})_{OLS} &= \|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|_2 \\ &= \boldsymbol{\varepsilon}^\top \boldsymbol{\varepsilon} \\ &= (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta})^\top (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}) \\ &= (\mathbf{Y}^\top - \mathbf{X}^\top \boldsymbol{\beta}^\top) (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}) \\ &= \mathbf{Y}^\top \mathbf{Y} - \mathbf{Y}^\top \mathbf{X}\boldsymbol{\beta} - \boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{Y} + \boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} \\ &= \mathbf{Y}^\top \mathbf{Y} - (\boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{Y})^\top - \boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{Y} + \boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} \\ &= \mathbf{Y}^\top \mathbf{Y} - 2(\boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{Y}) + \boldsymbol{\beta}^\top \mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} \end{aligned} \quad (3)$$

Next, a partial derivative of β is performed to obtain the minimum value of the equation:

$$\begin{aligned}\frac{\partial(\mathbf{S}(\beta)_{OLS})}{\partial(\beta)} &= \frac{\partial(\mathbf{Y}^\top \mathbf{Y} - 2(\beta^\top \mathbf{X}^\top \mathbf{Y}) + \beta^\top \mathbf{X}^\top \mathbf{X} \beta)}{\partial(\beta)} \\ &= 0 - 2\mathbf{X}^\top \mathbf{Y} + \mathbf{X}^\top \mathbf{X} \beta + (\beta^\top \mathbf{X}^\top \mathbf{X})^\top \\ &= 2\mathbf{X}^\top \mathbf{Y} + \mathbf{X}^\top \mathbf{X} \beta + \mathbf{X}^\top \mathbf{X} \beta\end{aligned}\quad (4)$$

and then equating it to zero, we get:

$$\begin{aligned}0 &= 2\mathbf{X}^\top \mathbf{Y} + \mathbf{X}^\top \mathbf{X} \beta + \mathbf{X}^\top \mathbf{X} \beta \\ 2\mathbf{X}^\top \mathbf{X} \beta &= 2\mathbf{X}^\top \mathbf{Y} \\ \mathbf{X}^\top \mathbf{X} \beta &= \mathbf{X}^\top \mathbf{Y} \\ \hat{\beta}_{OLS} &= (\mathbf{X}^\top \mathbf{X})^{-1}(\mathbf{X}^\top \mathbf{Y})\end{aligned}\quad (5)$$

Therefore, Equation (5) as the solution of the OLS method.

2.2. Ridge Regression. Ridge regression is the result of the least squares method with the addition of a bias value c to the correlation matrix and the variables are transformed using the centering and scaling method, the selection of the bias constant c is a very instrumental thing in Ridge regression [10]. The penalty in ridge with the following constraints:

$$\|\beta\|_2 \leq t, t > 0$$

Loss function ridge regression is as follow:

$$\begin{aligned}\mathbf{S}(\beta)_R &= \|\mathbf{Y} - \mathbf{X}\beta\|_2 + c\|\beta\|_2 \\ &= \epsilon^\top \epsilon + c\beta^\top \beta \\ &= (\mathbf{Y} - \mathbf{X}\beta)^\top (\mathbf{Y} - \mathbf{X}\beta) + c\beta^\top \beta \\ &= (\mathbf{Y}^\top - \mathbf{X}^\top \beta^\top)(\mathbf{Y} - \mathbf{X}\beta) + c\beta^\top \beta \\ &= \mathbf{Y}^\top \mathbf{Y} - \mathbf{X}^\top \beta^\top \mathbf{Y} - \mathbf{Y}^\top \mathbf{X} \beta + \beta^\top \mathbf{X}^\top \mathbf{X} \beta + c\beta^\top \beta \\ &= \mathbf{Y}^\top \mathbf{Y} - 2\beta^\top \mathbf{X}^\top \mathbf{Y} + \beta^\top \mathbf{X}^\top \mathbf{X} \beta + c\beta^\top \beta\end{aligned}\quad (6)$$

Next, a partial derivative of β is performed to obtain the minimum value of the equation:

$$\begin{aligned}\frac{\partial(\mathbf{S}(\beta)_R)}{\partial(\beta)} &= \frac{\partial(\mathbf{Y}^\top \mathbf{Y} - 2\beta^\top \mathbf{X}^\top \mathbf{Y} + \beta^\top \mathbf{X}^\top \mathbf{X} \beta + c\beta^\top \beta)}{\partial(\beta)} \\ &= 0 - 2\mathbf{X}^\top \mathbf{Y} + \mathbf{X}^\top \mathbf{X} \beta + (\beta^\top \mathbf{X}^\top \mathbf{X})^\top + 2c\beta \\ &= -2\mathbf{X}^\top \mathbf{Y} + \mathbf{X}^\top \mathbf{X} \beta + \mathbf{X}^\top \mathbf{X} \beta + 2c\beta \\ &= -2\mathbf{X}^\top \mathbf{Y} + 2\mathbf{X}^\top \mathbf{X} \beta + 2c\beta\end{aligned}\quad (7)$$

and then equating it to zero, we get:

$$\begin{aligned}
0 &= -2\mathbf{X}^\top \mathbf{Y} + 2\mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} + 2c\boldsymbol{\beta} \\
\mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} + c\boldsymbol{\beta} &= \mathbf{X}^\top \mathbf{Y} \\
(\mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} + c\mathbf{I})\boldsymbol{\beta} &= \mathbf{X}^\top \mathbf{Y} \\
\hat{\boldsymbol{\beta}}_R &= (\mathbf{X}^\top \mathbf{X}\boldsymbol{\beta} + c\mathbf{I})^{-1}(\mathbf{X}^\top \mathbf{Y})
\end{aligned} \tag{8}$$

Therefore, Equation (8) as the solution of the ridge regression.

2.3. LASSO Regression. LASSO (Regression Least Absolute Shrinkage and Selection Operator) is one of the shrinkage methods to overcome multicollinearity problems. The LASSO method is a method introduced by Tibshirani in 1996 [11] after the LAR (Least Angle Regression) method introduced by Effron in 2004 by changing the penalty in Ridge regression in L_1 regularization. This regularization is used to reduce overfitting by adding L_1 and L_2 penalty factors where L_1 regularization is called LASSO regression which uses L_1 penalty, an approach that penalizes the absolute size of the coefficients. Whereas L_2 regularization is called Ridge regression which uses an L_2 penalty, which is an approach that penalizes the squared size of the coefficients. LASSO aims to improve the estimation of simple linear regression. The penalty in LASSO with the following constraints:

$$\|\boldsymbol{\beta}\|_1 \leq t, t > 0$$

. The value of t above is a quantity that checks the amount of shrinkage in the LASSO coefficient estimates where $t \geq 0$. If the estimator $\hat{\boldsymbol{\beta}}$ is a least squares estimator and $t_0 = \|\boldsymbol{\beta}\|_1$, then values of $t < t_0$ will lead to solving classical regression with OLS estimators that shrink towards zero, and allow some coefficients to also shrink exactly towards zero. Loss function for LASSO regression is as follow:

$$\mathbf{S}(\boldsymbol{\beta})_{LASSO} = \|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|_2 + \alpha\|\boldsymbol{\beta}\|_1 \tag{9}$$

Unlike OLS estimation of linear regression in the equation (4-5) and ridge regression in the equation (7-8), LASSO regression cannot find direct results for the beta derivative so that one way that can be done is to find the iteration value that minimizes the loss function. The coefficient estimates in LASSO regression are written as follows [11]:

$$\hat{\boldsymbol{\beta}}_{LASSO} = \arg \min_{\boldsymbol{\beta}} (\|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|_2 + \alpha\|\boldsymbol{\beta}\|_1) \tag{10}$$

2.4. Elastic Net Regression. Elastic net is a penalty regression method similar to ridge regression and LASSO that can overcome the problem of multicollinearity assumption 2005 [12]. Elastic net combines the penalty between Ridge regression and LASSO. Elastic net can overcome the problem of high correlation and has the properties of variable selection and shrinkage of the estimation coefficient. Zou and Hastie (2005) [12] introduced the Elastic Net penalty as follows:

$$(1 - \lambda)\|\boldsymbol{\beta}\|_2 + \lambda\|\boldsymbol{\beta}\|_1 \leq t, t > 0$$

If $\lambda = 0$, the Elastic Net regression becomes a Ridge regression, while if $\lambda = 1$, the Elastic Net becomes a LASSO penalty. Elastic Net regularization has a shrinkage of the coefficient of correlated predictor variables like Ridge and LASSO. Loss function for Elastic Net regression is as follow:

$$S(\beta)_{ELN} = \|\mathbf{Y} - \mathbf{X}\beta\|_2 + \gamma ((1 - \lambda)\|\beta\|_2 + \lambda\|\beta\|_1) \quad (11)$$

The coefficient estimates in Elastic Net regression are written as follows:

$$\hat{\beta}_{ELN} = \arg \min_{\beta} (\|\mathbf{Y} - \mathbf{X}\beta\|_2 + \gamma ((1 - \lambda)\|\beta\|_2 + \lambda\|\beta\|_1)) \quad (12)$$

2.5. Artificial Neural Network (ANN). Artificial Neural Network (ANN) is an algorithm that has the same structure as the performance of the human brain in learning the pattern of data [13]. In an ANN cell, the weights and biases function as variables for the input data to be output. The weight and bias values are determined by the backpropagation method, which is an adjustment process so that the prediction results are close to the original value. This algorithm works by doing a back pass for each forward pass while adjusting the weights and biases. The process is assisted by an algorithm known as optimizers.

2.6. Deep Neural Network (DNN). Neural networks with multiple hidden layers are called deep neural networks (DNN) and the practice of training those networks are referred to as deep learning. Deep neural networks trained to adaptive to varied number of levels and nodes at each level, performance complex tasks, modeling the multiple outcomes.

3. MATERIAL AND METHOD

3.1. Variable. This paper uses data taken from BPS in 2018. The description of each variable can be explained below:

- Y : Number of villages with gambling occurrences in the last year by province, 2018
- X_1 : Open unemployment rate (TPT) by province in 2018
- X_2 : Education completion rate by education level (SMA) and province, 2018
- X_3 : Average monthly expenditure per capita on food in urban and rural areas by province (IDR), 2018
- X_4 : Average monthly non-food expenditure per capita in urban and rural areas by province (IDR), 2018
- X_5 : Average hourly wage of workers by province (IDR/hour)

The number of observations is 35 (35 provinces) with the dependent variable (Y) being the number of gambling cases in each province. This study uses log transformation to the dependent variable with the statistics descriptive of the variables is shown in **Table 1** below:

Variables	Minimum	Maximum	Mean
Y	35.0	1947.0	376.8
$\ln(Y)$	3.555	7.574	5.422
X_1	1.400	8.470	4.803
X_2	29.56	83.48	61.19
X_3	402922	847847	565941
X_4	301832	1191310	576476
X_5	11359	25987	15911

TABLE 1. Statistics Descriptive

3.2. Method. This paper uses natural logarithm ($\ln$) transformation on the dependent variable ($y^* = \ln(y)$) to increase the R-square value of the regression model. After the transformation, modeling is done using several regression methods such as linear, ridge, lasso, and elastic net. From each model, the error is calculated using the following formula:

$$\varepsilon = y^* - \hat{y}_{\text{single}}^* \quad (13)$$

The error is then used as input to the Deep Neural Network (DNN) with the aim of filling the gap between the original value and the predicted value and produce the estimated error ($\hat{\varepsilon}$). The final prediction of this hybrid model is formulated as:

$$\hat{y}_{\text{final}}^* = \hat{y}_{\text{single}}^* + \hat{\varepsilon}. \quad (14)$$

To evaluate the performance of the model, Root Mean Square Error (RMSE) and Mean Absolute Percentage Error (MAPE) is used as a comparison metric with the following formula:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i^* - \hat{y}_{\text{final}}^*)^2} \quad (15)$$

$$MAPE = \frac{1}{n} \sum_{i=1}^n \left| \frac{y_i^* - \hat{y}_{\text{final}}^*}{y_i^*} \right| \times 100\% \quad (16)$$

4. RESULTS AND DISCUSSION

4.1. Variable Selection. In this subsection, the first step is the selection of independent variables using the backward method in multiple linear regression. This process starts by including all five independent variables along with their intercepts into the initial model. Next, the variables are selected one by one, gradually removing the variables that have the least influence on the model. This process continues until only those variables remain that have a significant contribution to the dependent variable. The final results of this selection show that the three selected independent variables are X_2 , X_3 , and X_4 .

Variables in Model	Description
intercept, X_1, X_2, X_3, X_4, X_5	$\beta_1, \beta_2, \beta_3, \beta_5$ not significant β_0, β_4 significant
X_1, X_2, X_3, X_4, X_5	β_1, β_5 not significant $\beta_2, \beta_3, \beta_4$ significant
X_2, X_3, X_4, X_5	$\beta_2, \beta_3, \beta_4$ not significant β_5 significant
X_2, X_3, X_4	$\beta_2, \beta_3, \beta_4$ significant

TABLE 2. Variable Selection

4.2. Hybrid Model. In this subsection, we performed modeling using several single models, namely linear, ridge, lasso, and elastic net models. Each of these models produces an error which is then used as input for the Deep Neural Network (DNN). To improve the performance of the DNN, we performed hyperparameter tuning using the grid search method with a range of 1:50 for each layer on three different layers. **Table 3** shows the best results of the number of neurons in each layer obtained from the tuning process with each RMSE.

Model	Layer 1	Layer 2	Layer 3	RMSE
Linear	9	19	1	0.901931
Ridge	20	5	14	0.749255
Lasso	18	13	18	0.611645
Elastic Net	17	7	6	0.764042

TABLE 3. Best of Hyperparameter for DNN

The results for various combination are given in **Table 4**. Based on this table, it can be seen that the RMSE by the hybrid model is smaller than that of single models such as linear, ridge, lasso, and elastic net models. This shows that the combination of several regression models with the use of Deep Neural Network (DNN) as the final stage of modeling is able to provide more accurate predictions.

Model	RMSE	MAPE	RMSE Hybrid	MAPE Hybrid
Linear	1.03137	15.928%	0.86786	13.347%
Ridge	0.81773	13.144%	0.81446	13.091%
Lasso	0.86023	13.611%	0.56669	7.2017%
Elastic Net	0.86377	13.730%	0.71794	11.219%

TABLE 4. Model Performance

This hybrid approach strengthens the prediction results by utilizing the advantages of each regression model and DNN shown by **Figure 1**, thus capturing complex patterns that a single model may not be able to identify effectively.

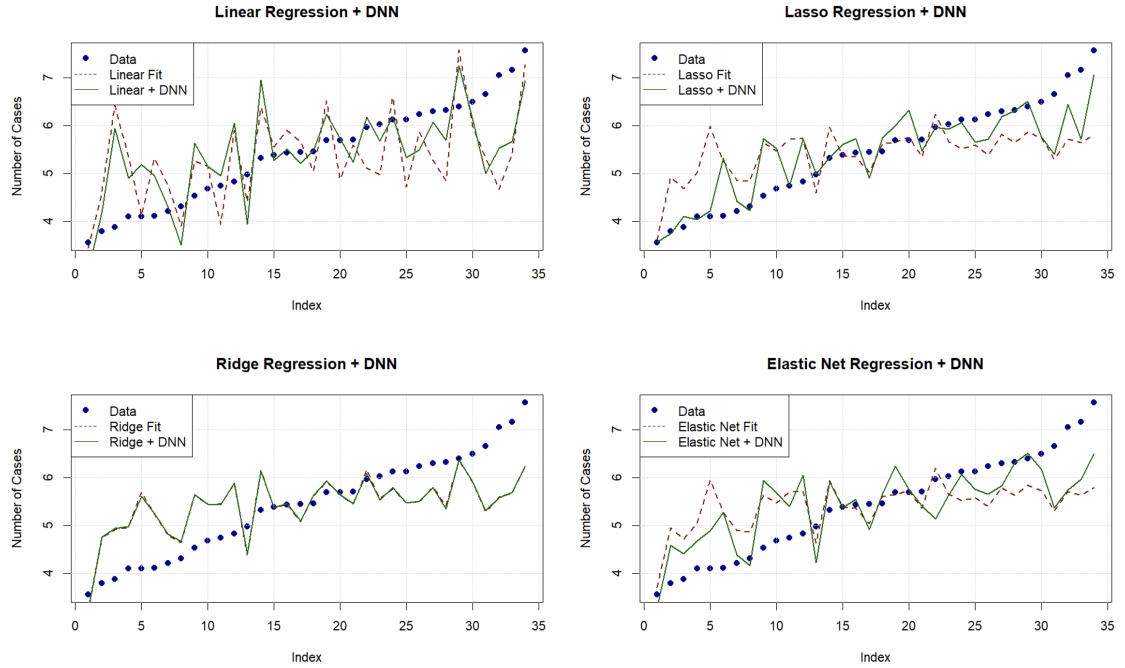


FIGURE 1. Comparison of Single Model vs Hybrid Model

5. CONCLUSION

The conclusion of this study shows that the use of linear regression and its extensions such as ridge, lasso, and elastic net can be maximized by combining it with a deep learning network. This hybrid approach is proven to produce smaller errors compared to a single model, indicating that this combination is able to provide more accurate and efficient predictions. Thus, hybrid models that integrate linear regression and deep learning networks offer a more robust solution in handling data complexity and improving predictive modeling performance.

REFERENCES

- [1] D. Khoerunisa, I. Nurahmadi, J. A. Sari, S. Wianti, and Y. E. Y. Siregar, "Judi online sebagai faktor penyebab permasalahan perceraian di kabupaten bekasi:(studi kasus pada kecamatan cikarang utara kabupaten bekasi)," *Kultura: Jurnal Ilmu Hukum, Sosial, dan Humaniora*, vol. 2, no. 2, pp. 63–70, 2024. [📄](#)

- [2] L. Wahkidi, E. Puspitasari, and T. Tamrin, "Hubungan tingkat kecanduan dengan tingkat kecemasan pelaku judi online di wilayah kecamatan toroh," *Jurnal Ilmu Keperawatan Komunitas*, vol. 5, no. 2, pp. 68–76, 2022. [\[1\]](#)
- [3] I. T. Jadidah, U. M. Lestari, K. A. S. Fatiha, R. Riyani, and Wulandari, "Analisis maraknya judi online di masyarakat," *Jurnal Ilmu Sosial dan Budaya Indonesia*, vol. 1, no. 1, pp. 20–27, 2023. [\[1\]](#)
- [4] R. Walpole, R. Myers, S. Myers, and K. Ye, *Probability & Statistics for Engineers & Scientists*. 1995. [\[1\]](#)
- [5] J. Y. L. Chan, S. M. H. Leow, K. T. Bea, W. K. Cheng, S. W. Phoong, Z. W. Hong, and Y. L. Chen, "Mitigating the multicollinearity problem and its machine learning approach: a review," *Mathematics*, vol. 10, no. 8, p. 1283, 2022. [\[1\]](#)
- [6] Z. Zhang, Z. Lai, Y. Xu, L. Shao, J. Wu, and G. S. Xie, "Discriminative elastic-net regularized linear regression," *IEEE Transactions on Image Processing*, vol. 26, no. 3, pp. 1466–1481. [\[1\]](#)
- [7] J. Ludbrook, "Linear regression analysis for comparing two measurers or methods of measurement: but which regression?," *Clinical and Experimental Pharmacology and Physiology*, vol. 37, no. 7, pp. 692–699, 2010. [\[1\]](#)
- [8] M. Firdaus, *Ekonometrika Suatu Pendekatan Aplikatif*. 2004. [\[1\]](#)
- [9] Weisstein and E. W., "vector norm", [mathworld.wolfram.com](https://mathworld.wolfram.com/vector-norm.html).
- [10] G. Azzahra, R. Herryanto, and F. Agustina, "Regresi ridge parsial untuk data yang mengandung masalah multikolinearitas," *Jurnal EurekaMatika*, vol. 8, no. 1, pp. 39–55, 2020. [\[1\]](#)
- [11] R. Tibshirani, "Regression shrinkage and selection via the lasso," *Journal of The Royal Statistical Society*, vol. 58, no. 1, pp. 267–288, 1996. [\[1\]](#)
- [12] H. Zou and T. Hastie, "Regularization and variable selection via the elastic net.," *Journal of the Royal Statistical Society Series B: Statistical Methodology*, vol. 67, no. 2, pp. 301–320. [\[1\]](#)
- [13] W. S. McCulloch and W. Pitts, "A logical calculus of the ideas immanent in nervous activity," *The bulletin of mathematical biophysics*, vol. 5, no. 7, pp. 115–133, 1943. [\[1\]](#)