

Psychologically Informed Instagram Marketing Analytics Pipeline Using Funnel Metrics and Multi-Criteria Decision Analysis: A Daycare Case Study

Hardika Khusnuliawati¹, Anniez Rachmawati Musslifah², Rusnandari Retno Cahyani³

Departments of Informatics, Universitas Sahid Surakarta, Jl. Adi Sucipto, Surakarta 57145, Central Java, Indonesia¹

Departments of Psychology, Universitas Sahid Surakarta, Jl. Adi Sucipto, Surakarta 57145, Central Java, Indonesia²

Departments of Business Administration, Universitas Sahid Surakarta, Jl. Adi Sucipto, Surakarta 57145, Central Java, Indonesia³

✉ Corresponding Author's Email: khusnuliawati@usahidsolo.ac.id

INFO ARTIKEL

History Article:

Submitted 08 22, 2026
Received 08 26, 2026
Published 08 27, 2026

ABSTRACT

Micro-scale social media accounts often lack sufficient observations for predictive analytics but still require defensible content and campaign decisions. This design-science proof of concept develops an Instagram marketing analytics pipeline combining funnel ratios, CRITIC weighting, TOPSIS ranking, and full-range weight sensitivity analysis (FRWSA). The pipeline was applied to ten complete-case organic posts and three advertising campaigns from an Indonesian Islamic Montessori daycare account. For organic posts, CRITIC assigned weights of 0.2764 to reach efficiency, 0.3887 to engagement rate, and 0.3349 to follow conversion. TOPSIS ranked the Thursday 04:02 post first with a closeness coefficient of 0.7503, while FRWSA found it dominant in 59.74% of 231 weight vectors. For advertising campaigns, equal weights were used because correlation-based weighting was unsuitable for only three alternatives. The 22–27 November 2025 campaign achieved a TOPSIS score of 1.0000 and ranked first across all 1,771 weight vectors, reflecting Pareto dominance rather than causality. Psychological and consumer-behavior theory informs the ordering of criteria by behavioral commitment, but no psychological state is measured or inferred. The study contributes a transparent decision-support pipeline for data-scarce social media accounts.

Keywords: Instagram Marketing Analytics; Funnel Metrics; CRITIC; TOPSIS; Weight Sensitivity Analysis; Daycare Case Study.

How to Cite:

Khusnuliawati, H., Musslifah, A. R., & Cahyani, R. R. (2026). Psychologically Informed Instagram Marketing Analytics Pipeline Using Funnel Metrics and Multi-Criteria Decision Analysis: A Daycare Case Study. LANCAH: Jurnal Inovasi Dan Tren, 4(2), 239-256. <https://journal.lembagakita.com/ljit/article/view/8504>

INTRODUCTION

Social media is the primary communication channel for small early childhood education providers, which typically operate without dedicated marketing staff or paid analytics tools. Instagram is especially prominent in this segment because visual storytelling communicates care quality, facilities and daily activities more persuasively than text. Mahatmi et al. (2022) report that using Instagram as a promotional medium supports audience engagement and brand awareness for a daycare provider, and Albastho et al. (2024) find that Instagram promotion combined with service quality is associated with parents' purchase decisions for childcare services. Related work in institutional settings shows that platform-native analytics can be mined for insight into audience response (Betri et al., 2024; Dewi et al., 2023), and that marketing strategy benefits from an explicit consumer-psychology perspective (Artura et al., 2024).

The practical problem, however, is not a shortage of numbers but a shortage of structure. A daycare account owner can read views, accounts reached, interactions and new follows for every post, and spend, clicks, messaging conversations and follows for every advertising campaign, yet these indicators point in different directions: the post with the widest exposure is rarely the post with the highest interaction intensity, and the cheapest click is not necessarily the click that produces a new follower. Social media measurement research is explicit that no single indicator suffices. Töllinen and Karjaluoto (2011) state that there is no universally applicable way to measure marketing performance and no broadly accepted single key performance indicator, and Trunfio and Rossi (2021) report that there is no solid consensus among scholars about measuring engagement on social media.

This article therefore treats content evaluation as a multi-criteria decision problem rather than a prediction problem. Multi-criteria decision analysis is designed exactly for situations in which several conflicting indicators must be combined into a defensible ordering, and it has been applied to digital marketing decisions such as platform selection using CRITIC weighting with TOPSIS ranking (Wegner et al., 2023). Two design constraints shape our approach. First, weights cannot be imported from the literature in multi-attribute value theory a weight is a scaling constant whose magnitude must follow the range of outcomes in the local decision context (Fischer, 1995; von Nitzsch and Weber, 1993), and the meaning of a weight depends on the measurement units and the aggregation rule adopted (Wedley et al., 1999). Second, because any single weight vector is contestable, the ordering it produces must be stress-tested rather than asserted; Gaona et al. (2025) show that TOPSIS is sensitive to weight assignment and propose enumerating the whole discretised weight simplex, reporting robustness as a dominance frequency.

The study is framed as a design-science proof of concept and an exploratory single-case study. Its object is the pipeline: a reproducible sequence that ingests platform-native counts, derives funnel ratios, applies objective weighting where the data permit it, ranks alternatives, and reports the stability of that ranking honestly. The empirical material, one Indonesian daycare account with ten complete-case organic posts and three completed campaigns, is used to demonstrate that the pipeline runs end to end and to illustrate how its output should be read. It is not used to establish an optimal posting time, a weekly pattern, or a causal effect of any design choice on audience behavior.

Three contributions follow. First, we give an operational definition of an Instagram marketing funnel for micro-accounts, in which every criterion is a ratio normalized by views so that posts of very different exposure remain comparable. Second, we show how the weighting problem can be handled defensibly at two very different sample sizes: objective CRITIC weights for the ten-post organic matrix,

and explicit equal weights as a neutral benchmark for the three-campaign matrix, where correlation-based weighting is not credible. Third, we demonstrate full-range weight sensitivity analysis as a mandatory reporting step, which converts a single ranking into a statement about how much of the weight space supports it. A psychologically informed reading is part of the design rather than a promise of future work: established behavioral theory organizes the criteria by increasing behavioral commitment and disciplines their interpretation, while the explicit non-inference rule keeps digital traces from being read as latent mental states.

LITERATURE REVIEW

Instagram Marketing for Early Childhood Education Providers

Studies of daycare and childcare marketing in Indonesia consistently describe Instagram as the main promotional instrument. Mahatmi et al. (2022) document a service-learning intervention in which Instagram content planning improved audience reach and brand recognition for a daycare, while Albastho et al. (2024) find a statistical association between Instagram promotion, perceived service quality and parents' decisions to purchase childcare services. Analytical work on institutional accounts (Betri et al., 2024; Dewi et al., 2023) shows that platform-native insight data can be treated as a legitimate research object. What this literature does not provide is a decision rule: it establishes that Instagram matters without specifying how an owner should compare posts that differ on several indicators at once.

Funnel Metrics and the Absence of Standard Weights

Marketing communication research organizes social media indicators along a funnel. Töllinen and Karjaluo (2011) reproduce a five-stage structure of exposure, influence, engagement, action or conversion, and retention, and argue that any single silver metric is inadequate. Pencarelli and Mele (2019) group social media metrics into activity, sharing or brand visibility, dimensional, engagement or interaction, and performance or business categories, note that only the performance category supports economic assessment, and leave the prioritisation of targets to each firm. Trunfio and Rossi (2021) show that a substantial share of engagement studies use normalized indexes, that is, interactions expressed relative to the audience that received the content, and record that although some studies weight interaction types differently, no transferable numeric weights are reported. Gräve (2019) provides the only relative importance figures we located for social media indicators, and they concern influencer attributes rather than post-level ratios. The literature therefore supports the choice and the ordinal hierarchy of criteria, but not their numeric weights.

Two recent studies strengthen the case for treating funnel stages as distinct measurement layers rather than as a single performance construct. Qin et al. (2024) operationalise brand awareness, consideration and revenue as separate funnel outcomes for 122,000 brands on an online marketplace and report that the same advertising or retail action can move one funnel stage while doing little for another, with the pattern differing by brand size and product category; their evidence argues against collapsing indicators into one interchangeable key metric, although their causal estimates and recommended budget allocations belong to a retail-media environment and are not transferable to a micro-scale Instagram account. Crowe et al. (2025) construct a funnel representation directly from observed behavioral engagement in a platform advertising dataset of more than 200 million interactions

and show that early-stage behaviors such as impressions and post engagement are qualitatively different from later actions such as landing-page clicks and conversions, and that the choice of metrics materially changes the funnel that is recovered. Both studies support the behavioral mapping of metrics to funnel depth adopted here; neither licenses any transfer of their coefficients, transition probabilities or scale-dependent conclusions to the present case, whose data are aggregate, cross-sectional and orders of magnitude smaller.

Multi-Criteria Decision Analysis in Digital Marketing

Multi-criteria decision analysis has been applied to digital marketing problems including social media platform selection with CRITIC weighting and TOPSIS ranking (Wegner et al., 2023), and to Instagram content ranking with subjectively assigned weights. In the latter case the authors themselves note that subjective assignment may introduce bias and reduce generalisability, which is precisely the weakness the present design avoids. Şahin (2020) classifies weighting methods as subjective, objective or hybrid, observes that attributes are legitimately determined through expert knowledge and literature review, and notes that CRITIC, entropy and standard deviation methods treat the decision matrix as the only source of information.

Objective Weighting with CRITIC and Its Preconditions

The CRITIC method derives weights from the contrast intensity of each criterion and from the conflict between criteria, measured through inter-criterion correlation (Diakoulaki et al., 1995). Its authors warn that incorporating several interdependent criteria can yield misleading results, and that different elicitation methods and different decision makers produce different weights for the same problem, which is the motivation for a data-driven alternative. Two preconditions matter here. CRITIC requires a complete numerical matrix and stable inter-criterion correlations (Wolny, 2025), and correlation estimates are inaccurate in small samples, stabilising only at sample sizes far larger than those available in micro-account analytics (Schönbrodt and Perugini, 2013). Entropy weighting is additionally distorted when the matrix contains zero values (Zhu et al., 2020). These constraints determine our two-track weighting design and the treatment of the anomalous zero-reach record described in Section 3.5.

Weight Sensitivity and Robust Reporting

Gaona et al. (2025) integrate TOPSIS with a full-range weight sensitivity analysis: because incorrect or arbitrary weights can compromise the reliability of the decision process, they enumerate all admissible weight vectors on a discretised simplex in steps of 0.05 and summarise robustness as dominance frequency, the share of the simplex in which an alternative ranks first. They also use a fixed set of equal weights as a neutral benchmark for its simplicity. Both devices are adopted here, and dominance frequency becomes the principal robustness statistic reported in Section 4.

A Psychologically Informed Reading of Observable Funnel Behavior

The term psychologically informed is used here in a precise and deliberately modest sense: established psychological and consumer-behavior theory supplies the organising lens that distinguishes the observable actions recorded by the platform and orders them by increasing behavioural commitment, from an opportunity for exposure and attention, through low-cost interaction, to purposive action such

as a link click, a message or a new follow. Funnel models in marketing communication already encode this progression (Töllinen and Karjaluoto, 2011; Pencarelli and Mele, 2019), and recent funnel analytics confirms that actions at different depths behave as distinct measurement layers rather than as interchangeable expressions of one construct (Qin et al., 2024; Crowe et al., 2025). Behavioural commitment is therefore an ordering property of the actions themselves, not a claim about what the actor felt or intended.

Psychological research simultaneously imposes the restraint that gives the lens its discipline. Attitude and affect cannot be read off a behavioural trace unless action, target, context and time are specified at the same level of generality, a correspondence condition that aggregate platform counts do not satisfy: the export records that an action occurred, not who performed it, towards what object, in what situation or with what deliberation. Where psychological constructs are studied directly, they are measured with dedicated validated instruments (Kaur et al., 2016), and engagement itself is conceptualised across behavioural, affective and cognitive dimensions of which platform indexes capture only the behavioural one (Trunfio and Rossi, 2021). Documented diurnal and weekly rhythms in expressed affect and online activity (Golder and Macy, 2011; Sano et al., 2019) and day-to-day variation in parenting stressors (Hibel et al., 2012) make it plausible that caregiver receptivity varies, but they license no inference from a count to a state. Accordingly, reach is treated as exposure opportunity and not as awareness, interactions as low-cost behavioural response and not as emotion, and messages or follows as purposive behavioural signals at a deeper funnel position and not as measured intention or trust.

This position has two concrete consequences for the present design rather than for a future study. First, it governs the grouping and interpretation of the criteria: each funnel stage contributes one criterion at the depth at which it is actually observed, ratios are normalised by views so that alternatives of unequal exposure remain comparable within a stage, and the organic and advertising matrices are kept apart because they observe different depths. Second, it imposes restraint in reporting: heterogeneous indicators are not collapsed into a single performance construct outside the explicit aggregation rule, and no psychological label is attached to a closeness coefficient, a weight or a dominance frequency.

METHODS OF RESEARCH

Research Design

The study follows a design-science logic: the artefact under study is an analytics pipeline, and the empirical material serves as a demonstration case. The design is exploratory, descriptive and non-experimental. No hypothesis is tested, no inferential statistic is reported, and no causal claim is made. Two decision problems are addressed separately: ranking organic posts on funnel efficiency, and ranking completed advertising campaigns on objective-aligned efficiency. In both cases the deliverable is an ordering accompanied by an explicit statement of how much of the weight space supports it.

Figure 1 summarises the five-stage analytical pipeline and the non-numeric role of the psychology-informed interpretive lens. The interpretive lens is not an input variable and enters no computation: it organises the criteria by increasing behavioural commitment and disciplines the vocabulary in which the resulting ranking is read, while every quantity in Stages 2 to 5 is derived from observed platform metrics alone.

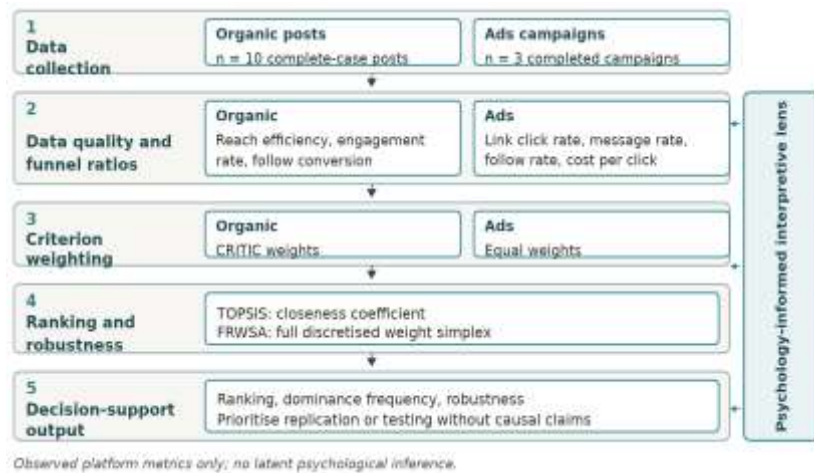


Figure 1. Research Flow of the Psychology-Informed Funnel-Metrics and Multi-Criteria Decision-Analysis Pipeline
(Source: Based on the Author's Own Analysis)

Case Context and Data Sources

The case is a single Instagram business account operated by an Islamic Montessori daycare in Central Java, Indonesia. Two exports were obtained from the account owner with permission. The organic export covers eleven posts published between 14 May and 29 July 2025 with views, accounts reached, interactions and new follows per post, together with publication date, day and time. The advertising export summarises three completed campaigns run between 22 November and 11 December 2025, all with the objective of obtaining more profile visits, with budget, spend, views, viewers, three-second plays, link clicks, reactions, comments, shares, saves, messaging conversations and new follows. Both exports are platform-native summaries; the platform documents that such metrics are estimated and that some are still in development (Meta Platforms, n.d.-a, n.d.-b).

Funnel Stages and Metric Definitions

Four funnel stages are defined for the organic problem: exposure, measured by views and accounts reached; attention efficiency, measured by the share of views that corresponds to unique accounts reached; interaction, measured by total interactions; and acquisition, measured by new follows. For the paid problem a traffic stage is added, measured by link clicks, and a conversation stage, measured by messaging conversations, together with a cost stage measured in Indonesian rupiah per click. Because absolute counts scale with exposure, every criterion is a ratio normalised by views, following the normalised-index practice documented by Trunfio and Rossi (2021). For organic post i :

$$RE_i = Reach_i / Views_i \quad (1)$$

$$ER_i = Interactions_i / Views_i \quad (2)$$

$$FC_i = Follows_i / Views_i \quad (3)$$

Where RE is reach efficiency, ER is engagement rate and FC is followed conversion, and all three are benefit criteria (Equations 1 to 3). Engagement rate is reported as a proportion in the decision matrix and as a percentage in the text; it is defined on views, not on followers, and it is not interchangeable with reach-based or follower-based engagement definitions. For advertising campaign k :

$$LCR_k = Clicks_k / Views \quad (4)$$

$$MR_k = Messages_k / Views \quad (5)$$

$$FR_k = Follows_k / Views \quad (6)$$

$$CPC_k = Spend_k / Clicks \quad (7)$$

Where LCR is link click rate, MR is message rate and FR is followed rate, all benefit criteria, and CPC is cost per click, a cost criterion (Equations 4 to 7). Interactions in the organic export are the platform aggregate of likes, comments, saves and shares; in the advertising export the same aggregate is not reported directly, which is one reason the two matrices are kept apart.

The two exports are consequently mapped to different funnel depths, following the stage-specific measurement practice of Qin et al. (2024) and the behavior-based funnel construction of Crowe et al. (2025): the organic criteria reach only as far as exposure efficiency, interaction and account acquisition, whereas the advertising criteria additionally cover a traffic stage measured by link clicks, a conversation stage measured by messaging conversations and a cost stage measured in rupiah per click, because those variables are observed only in the advertising export. Consistent with both studies, the metrics are not treated as interchangeable: a ratio defined at one funnel depth is never substituted for a ratio defined at another, no composite single indicator is formed outside the explicit TOPSIS aggregation, and each criterion is interpreted only at the depth at which it is actually observed.

Why Organic and Paid Matrices are Not Pooled

The two decision problems are analysed as separate matrices for four reasons. The alternatives differ in kind: a single post versus a multi-day campaign with its own budget, duration and targeting. The exposure mechanism differs: organic distribution is algorithmic and unpaid, whereas paid distribution is purchased and optimised by the platform towards a chosen objective. The available criteria differ: cost per click and message rate exist only for campaigns, and accounts reached is available only for organic posts in these exports. Finally, pooling would make the ranking depend on an arbitrary exchange rate between paid and organic exposure, which no source in our reference set supplies. Pooling was therefore rejected as indefensible rather than merely inconvenient.

Treatment of the Zero-Reach Record

One organic record, published on 29 July 2025, reports accounts reached equal to zero together with 448 views and 36 interactions. The platform defines views as the number of times content was played or displayed and accounts reached as the number of unique accounts that saw the content at least once (Meta Platforms, n.d.-a). Since a display requires at least one account to have seen the content, the combination of zero reach with positive views is inconsistent with these definitions. We stress that no consulted platform page states explicitly that zero reach with positive views is impossible; the inconsistency is an inference from the official definitions, and the platform documentation notes that individual metrics may be returned as zero and that insight values are estimated or in development (Meta Platforms, n.d.-a, n.d.-b).

Accordingly, the value is treated as not available rather than as a valid zero, and the record is excluded from the primary reach-based decision matrix, which is therefore a complete-case matrix of ten posts. The raw zero remains disclosed in Table 1 so that no information is hidden. Two robustness handlings are reported qualitatively in Section 4.4: retaining the record with reach treated as a valid zero, which places it at the bottom of the reach efficiency criterion while giving it the highest engagement rate in the sample, and imputing its reach efficiency at the median of the remaining posts. Neither handling can be treated as authoritative, because both assign a value that the export does not contain; excluding the record from the reach-based matrix is the only handling that adds no unsupported information. Entropy weighting was rejected for the same reason: zero values distort its standardisation step (Zhu et al., 2020).

Criteria Selection and the Two-Track Weighting Decision

Criteria were selected from the literature, and weights were not. The funnel hierarchy of Töllinen and Karjaluo (2011) and the metric categories of Pencarelli and Mele (2019) justify including one exposure-efficiency criterion, one interaction criterion and one acquisition criterion for organic posts, and the campaign objective of obtaining more profile visits justifies the reduction of the nine available advertising indicators to a traffic criterion, a conversation criterion, an acquisition criterion and a single cost criterion. Numeric weights were not imported, because weights are scaling constants tied to the range of outcomes in the local decision context (Fischer, 1995; von Nitzsch and Weber, 1993) and their interpretation depends on units and aggregation rule (Wedley et al., 1999), and because the published multi-criteria studies in digital marketing weight different decision objects, such as platforms or strategies, rather than post-level funnel ratios.

For the organic matrix, weights are derived objectively with CRITIC from the local decision matrix itself, following the precedent set in digital marketing by Wegner et al. (2023) and the recommendation to use objective weighting when expert judgement is unavailable (Şahin, 2020). For the advertising matrix, neither CRITIC nor entropy is used: with three alternatives, inter-criterion correlations are not credible (Schönbrodt and Perugini, 2013) and interdependent criteria can mislead (Diakoulaki et al., 1995). Instead, the four selected criteria receive equal weights of 0.25 as an explicit neutral benchmark, exactly as Gaona et al. (2025) use fixed equal weights for simplicity and neutrality. In both tracks the reported weight vector is then stress-tested over the entire weight simplex.

CRITIC Weighting

The decision matrix is first standardised so that every benefit criterion is expressed on a common zero-to-one scale:

$$x'_{ij} = (x_{ij} - \min_j) / (\max_j - \min_j) \quad (8)$$

The information content of criterion j combines its contrast intensity, the standard deviation of the standardised column, with its conflict against the other criteria, expressed through the Pearson correlation r between standardised columns (Diakoulaki et al., 1995):

$$C_j = \sigma_j \cdot \sum_k (1 - r_{jk}) \quad (9)$$

$$w_j = C_j / \sum_j C_j \quad (10)$$

Equations 8 to 10 were implemented with the sample standard deviation. The resulting weights are reported in Table 2 together with the correlation sums, so that a reader can reproduce each step from the published matrix.

TOPSIS Ranking

Alternatives are ranked with TOPSIS. The decision matrix is first normalised by the vector norm of each column, then weighted:

$$n_{ij} = x_{ij} / \sqrt{\sum_i x_{ij}^2} \quad (11)$$

$$v_{ij} = w_j \cdot n_{ij} \quad (12)$$

The positive ideal solution takes the maximum weighted value on each benefit criterion and the minimum on each cost criterion, and the negative ideal solution takes the opposite; for the organic matrix all criteria are benefit criteria, whereas for the advertising matrix cost per click is a cost criterion. Separation measures and the closeness coefficient follow:

$$D_i^+ = \sqrt{\sum_j (v_{ij} - v_j^+)^2} \quad (13)$$

$$D_i^- = \sqrt{\sum_j (v_{ij} - v_j^-)^2} \quad (14)$$

$$CC_i = D_i^- / (D_i^+ + D_i^-) \quad (15)$$

Equations 11 to 15 yield a closeness coefficient bounded by zero and one, where a larger value

indicates a profile closer to the best observed performance in that specific matrix. The coefficient is a relative index within the sample and carries no absolute meaning: the lowest-ranked alternative necessarily receives a coefficient of zero when it defines the negative ideal on every criterion, which does not mean that it produced no value.

Full-Range Weight Sensitivity Analysis

Because the ranking depends on the weight vector, every ranking is recomputed over the entire admissible weight space, discretised in steps of 0.05 (Gaona et al., 2025). The admissible set is

$$W = \{w: w_j \in \{0, 0.05, \dots, 1\}, \sum_j w_j = 1\} \quad (16)$$

which contains 231 vectors for the three organic criteria and 1,771 vectors for the four advertising criteria. For each vector, TOPSIS is executed and the first-ranked alternative is recorded; ties share the credit equally. Dominance frequency is then

$$DF_i = (1 / |W|) \cdot \sum 1/[I \text{ rank first under } w] \quad (17)$$

Equations 16 and 17 turn a single ordering into a statement about the share of the weight space that supports it, which is reported alongside every ranking in Section 4. A dominance frequency of one means that the alternative would rank first under every admissible weight vector on the selected criteria, that is, it is Pareto dominant within that criteria set; a value close to one half means that the leading position is genuinely contestable.

Ethics, Reproducibility and the Computational Environment

The study analyses aggregate account-level metrics supplied by the account owner with permission for research use. No personal data of parents, children or staff were accessed, no individual user was identified or profiled, and no comment text, direct message or image of a child is reproduced. Post captions are referred to only in aggregate terms. All computations are deterministic and were performed in a spreadsheet workbook that stores every intermediate quantity, including the standardised matrix, the correlation matrix, the CRITIC information indices, the vector-normalised and weighted matrices, both ideal solutions, the closeness coefficients, and the complete enumeration of weight vectors with the winner of each. Because every step is elementary arithmetic, the results can be reproduced in any environment; we therefore state the workbook itself, rather than a software version number, as the reproducibility artefact, and we do not report library versions that we cannot verify.

RESULT AND DISCUSSION

Descriptive Funnel Profile of the Organic Posts

Table 1 reports the organic decision matrix. Across the ten complete-case posts, the account accumulated 38,581 views, 20,805 accounts reached, 592 interactions and 51 new follows, giving pooled values of 0.539 for reach efficiency, 1.53 per cent for engagement rate and 0.132 per cent for follow conversion. The three criteria separate the posts differently. Reach efficiency ranges from 0.173 to 0.760, engagement rate from 0.68 to 2.64 per cent, and follow conversion from 0 to 0.304 per cent. The post with the highest reach efficiency, the 23 July 18:12 post at 0.760, has the second lowest engagement rate, while the post with the highest engagement rate, the 16 May 03:55 post at 2.64 per cent, ranks fifth on reach efficiency. Follow conversion is the most skewed criterion: a single post, 15 May at 04:02, accounts for 35 of the 51 new follows in the complete-case window, and two posts produced none. This pattern of conflicting criteria is exactly the configuration for which multi-criteria aggregation is appropriate, and it is also what CRITIC exploits when it assigns weights.

Table 1. Organic Decision Matrix. Ratios Follow Equations 1 to 3. The Final Row is Excluded from the Primary Matrix Because the Reported Reach is Inconsistent with the Platform Definitions (Section 3.5)

No	Date	Day	Time	Views	Reach	Inter.	Follows	RE	ER (%)	FC (%)
1	14 May 2025	Wednesday	07:26	2,604	1,422	66	1	0.546	2.53	0.038
2	15 May 2025	Thursday	04:02	11,495	7,543	153	35	0.656	1.33	0.304
3	16 May 2025	Friday	03:55	2,727	1,650	72	5	0.605	2.64	0.183
4	17 May 2025	Saturday	21:12	2,556	1,604	65	4	0.628	2.54	0.156
5	26 May 2025	Monday	00:39	1,877	1,249	32	0	0.665	1.70	0.000
6	4 Jun 2025	Wednesday	19:04	2,864	1,535	58	2	0.536	2.03	0.070
7	4 Jun 2025	Wednesday	05:52	1,981	1,155	38	1	0.583	1.92	0.050
8	17 Jun 2025	Tuesday	23:03	1,396	915	29	2	0.655	2.08	0.143
9	23 Jul 2025	Wednesday	18:12	3,096	2,354	25	1	0.760	0.81	0.032
10	23 Jul 2025	Wednesday	02:02	7,985	1,378	54	0	0.173	0.68	0.000
Total	—	—	—	38,581	20,805	592	51	0.539	1.53	0.132
11*	29 Jul 2025	Tuesday	04:49	448	0*	36	0	n.a.	8.04	0.000

Note. RE = reach efficiency, ER = engagement rate, FC = follow conversion. Totals are pooled values over the ten complete cases. *Record excluded from the primary matrix; its reach is treated as not available rather than as a valid zero.

The excluded record, published on 29 July at 04:49, is shown in the final row of Table 1 for completeness: 448 views, zero accounts reached, 36 interactions and no new follows. Its nominal engagement rate of 8.04 per cent is more than three times the highest value among the complete-case posts and coexists with the lowest view count in the export, which reinforces the interpretation of the record as a reporting artefact rather than an exceptional creative success.

CRITIC Weights for the Organic Matrix

Table 2 reports the CRITIC computation. The standardised columns have sample standard deviations of 0.2680, 0.3564 and 0.3201 for reach efficiency, engagement rate and follow conversion respectively, and the inter-criterion correlations are moderate and positive: 0.280 between reach efficiency and engagement rate, 0.339 between reach efficiency and follow conversion, and 0.260 between engagement rate and follow conversion. The resulting information indices are 0.3701, 0.5203 and 0.4483, and the normalised weights are 0.2764 for reach efficiency, 0.3887 for engagement rate and 0.3349 for follow conversion.

The ordering of these weights is a property of this ten-post matrix, not a general finding. Engagement rate receives the largest weight because it combines the highest contrast with the weakest average correlation to the other two criteria; reach efficiency receives the smallest weight because its values are the most tightly clustered. A different observation window, or the inclusion of the excluded record, would change these numbers. That is why the weights are treated as one point in a weight space rather than as the correct weights, and why Section 4.4 reports the full enumeration.

Table 2. CRITIC Weighting of the Organic Criteria (Equations 8 To 10), Computed on the Ten Complete-Case Posts

Criterion	σ_j	$\Sigma(1-r)$	C_j	w_j
Reach efficiency	0.2680	1.3808	0.3701	0.2764
Engagement rate	0.3564	1.4599	0.5203	0.3887
Follow conversion	0.3201	1.4003	0.4483	0.3349
Sum	—	—	1.3387	1.0000

Note. Correlations between standardised criteria: reach efficiency and engagement rate 0.280; reach efficiency and follow conversion 0.339; engagement rate and follow conversion 0.260. All criteria are benefit criteria.

Organic CRITIC-TOPSIS Ranking

Table 3 reports the closeness coefficients obtained with Equations 11 to 15 under the CRITIC weights. The 15 May post published at 04:02 on Thursday ranks first with 0.7503, followed by the 16 May Friday post at 0.6719, the 17 May Saturday post at 0.6077 and the 17 June Tuesday post at 0.5464. The bottom of the ranking is occupied by the 23 July 02:02 event announcement, whose coefficient is exactly zero because it simultaneously defines the negative ideal on all three criteria: it has the lowest reach efficiency, the lowest engagement rate and no new follows.

Table 3. Organic CRITIC-TOPSIS Ranking of the Ten Complete-Case Posts (Equations 11 to 15)

Rank	Post	Day and time	CC_i
1	2	Thursday 04:02	0.7503
2	3	Friday 03:55	0.6719
3	4	Saturday 21:12	0.6077
4	8	Tuesday 23:03	0.5464
5	1	Wednesday 07:26	0.3859
6	6	Wednesday 19:04	0.3742
7	7	Wednesday 05:52	0.3401
8	5	Monday 00:39	0.2818
9	9	Wednesday 18:12	0.2693
10	10	Wednesday 02:02	0.0000

Note. CC = closeness coefficient, a relative index within this matrix only. Post numbers follow Table 1. Day and time are descriptive labels, not explanatory variables.

The leading post owes its position primarily to follow conversion, on which it is the best-performing alternative by a wide margin, and secondarily to reach efficiency; its engagement rate is in fact below the sample median. Conversely the two posts with the highest engagement rates rank second and third rather than first. The ranking must therefore not be read as a ranking of creative quality, and still less as evidence about publication time. Four of the ten posts were published on Wednesday and they occupy ranks 5, 6, 7 and 9, but with a single post on most other days no comparison across days of the week is possible.

Robustness of the Organic Ranking

Table 4 reports the full-range weight sensitivity analysis over the 231 admissible weight vectors. The leading post ranks first in 138 vectors, a dominance frequency of 59.74 per cent. The second-ranked post leads in 64 vectors, or 27.71 per cent, the third-ranked post in 15 vectors, or 6.49 per cent, and two further posts lead in a small number of vectors each: 8 vectors, or 3.46 per cent, for the 23 July 18:12 post, which is the reach efficiency leader, and 6 vectors, or 2.60 per cent, for the 17 June post. Five of the ten posts never rank first under any admissible weight vector.

Table 4. Full-Range Weight Sensitivity Analysis of the Organic Ranking Over All 231 Admissible Weight Vectors (Equations 16 and 17)

Post	Equivalent wins	Dominance freq. (%)	Robust rank
2	138	59.74	1
3	64	27.71	2
4	15	6.49	3
9	8	3.46	4
8	6	2.60	5
1, 5, 6, 7, 10	0	0.00	6 (tied)

Note. Step size 0.05 on the three-criterion simplex; ties receive fractional credit. Dominance frequency is the share of the weight space in which a post rank first.

The correct reading is that the top position is stable but not universal: a decision maker who weighted engagement rate heavily enough would prefer the 16 May post, and one who weighted reach efficiency almost exclusively would prefer the 23 July 18:12 post. What the enumeration does establish is a robust upper group, since the three leading posts together account for 93.94 per cent of the weight space, and a robust lower group of five posts that never lead. Regarding the excluded record, retaining it with reach treated as a valid zero would place it at the extreme of two criteria at once, giving it the maximum engagement rate and the minimum reach efficiency in the matrix, which inflates the contrast term of both criteria and changes the CRITIC weights; imputing its reach efficiency at the median of the remaining posts would instead invent an exposure value that the export does not report. Because both handlings alter the weighting inputs on the basis of an assumption rather than an observation, we report the complete-case matrix as the primary analysis and refrain from presenting either variant as a comparable ranking.

Advertising Campaigns: Criteria and Ranking

Table 5 reports the reduced advertising decision matrix. The three campaigns spent 83,931, 117,534 and 124,942 rupiah and generated 8,984, 15,657 and 13,597 views respectively. On the four objective-aligned criteria, the 22 to 27 November campaign attains the highest link click rate at 4.93 per cent, the highest message rate at 0.0445 per cent, the highest follow rate at 0.445 per cent and the lowest cost per click at 189 rupiah. The 24 November to 1 December campaign has the second-best link click rate and follow rate and the second-lowest cost per click, and the 6 to 11 December campaign has the worst link click rate, follow rate and cost per click, at 686 rupiah, while holding the second-highest message rate.

**Table 5. Reduced Advertising Decision Matrix Aligned with the Profile-Visit Campaign
Objective (Equations 4 to 7)**

Campaign	LCR (%)	MR (%)	FR (%)	CPC (IDR)
22–27 Nov 2025	4.931	0.0445	0.445	189
24 Nov–1 Dec 2025	2.280	0.0128	0.141	329
6–11 Dec 2025	1.339	0.0221	0.059	686

Note. LCR = link click rate, MR = message rate, FR = follow rate (all benefit); CPC = cost per click (cost). Spend and views: 83,931 and 8,984; 117,534 and 15,657; 124,942 and 13,597. Equal weights of 0.25 are used as a neutral benchmark because correlation-based weighting is not defensible at three alternatives.

Under equal weights of 0.25 the closeness coefficients are 1.0000, 0.3338 and 0.1226 respectively, as reported in Table 6. The leading campaign attains the maximum possible coefficient because it defines the positive ideal on three criteria and the negative ideal on none: its distance to the positive ideal is zero. This is an arithmetic consequence of dominance in a three-alternative matrix, not a measure of effect size.

Robustness of the Advertising Ranking

Table 6 also reports the full-range sensitivity analysis over the 1,771 admissible weight vectors of the four-criterion simplex. The 22 to 27 November campaign ranks first in all 1,771 vectors, a dominance frequency of 100 per cent, and neither of the other campaigns leads under any admissible weight vector. This result is entirely attributable to Pareto dominance: because the leading campaign is best on all three benefit criteria and cheapest on the cost criterion, no non-negative weight vector can reverse the comparison. The enumeration therefore adds no information beyond confirming dominance and is reported only for transparency.

**Table 6. Advertising TOPSIS Ranking Under Equal Weights and Full-Range Weight
Sensitivity Analysis Over All 1,771 Admissible Weight Vectors**

Campaign	CC _i	Rank	Dominance freq. (%)
22–27 Nov 2025	1.0000	1	100.00
24 Nov–1 Dec 2025	0.3338	2	0.00
6–11 Dec 2025	0.1226	3	0.00

Note. The leading campaign is Pareto dominant on the four selected criteria, so its dominance frequency of 100 per cent is a logical consequence and not causal evidence that its schedule, budget or creative material was superior.

It follows that the 100 per cent figure must not be read as strong evidence for a scheduling rule. The three campaigns differ simultaneously in budget, duration, creative material, audience targeting and calendar period, and all three were published in the afternoon or evening, so the design contains no contrast that would isolate any single factor. The observation that the dominant campaign happened to be launched on a Friday is a description of this sample and not evidence that Friday afternoons are preferable.

Discussion

What the Pipeline Contributes

The pipeline converts an unstructured export into an auditable decision record. Every step is explicit: the funnel stage each criterion represents, the normalisation that makes posts of different exposure comparable, the source of the weights, the ranking rule, and the share of the weight space that

supports the resulting order. For a provider with no analytics staff the value lies less in the specific ordering than in the discipline it imposes: it prevents selecting whichever indicator flatters a recent post, and a stakeholder who disputes the outcome must state a weight vector, which the enumeration in Table 4 evaluates immediately. The psychology-informed element of this contribution is equally methodological: it consists of a criterion architecture ordered by increasing behavioural commitment together with an explicit rule of non-inference from digital traces to latent mental states, and it should not be read as an empirical contribution to psychology.

Interpreting Robustness Correctly

The two robustness results illustrate opposite situations. In the organic matrix, a dominance frequency of 59.74 per cent means that the leading post is the modal winner rather than the unambiguous one, and that a defensible alternative preference exists for roughly two fifths of the weight space. Reporting the ranking without this figure would overstate the finding considerably. In the advertising matrix, a dominance frequency of 100 per cent is logically guaranteed by Pareto dominance and therefore says nothing about the reliability of the underlying measurements, which remain platform estimates from three uncontrolled campaigns. This asymmetry supports the recommendation of Gaona et al. (2025) that sensitivity analysis accompany every TOPSIS application, and it also shows that dominance frequency must be interpreted jointly with the structure of the matrix.

Funnel Implications for Organic and Paid Activity

Read as a funnel diagnosis rather than a scoreboard, the results are informative in a limited way. On the organic side, exposure efficiency and interaction intensity are not aligned in this account, and acquisition is concentrated in a single post, which suggests that follow-through to the account profile is the scarce step rather than exposure. On the paid side, the dominant campaign is the most efficient at every stage that the reduced criteria set represents, from clicks through messages to follows, while the campaign with the largest view volume, 15,657 views for 117,534 rupiah, converts that exposure least efficiently into clicks. A brand-awareness objective might value that larger upper-funnel volume, and our reduced criteria set deliberately does not: the campaigns were run with the objective of obtaining more profile visits, so the analysis is aligned with that objective and its conclusions do not extend to awareness goals. The pipeline is thus consistent with the stage-sensitive funnel analytics of Qin et al. (2024) and Crowe et al. (2025). Unlike Crowe et al. (2025), however, it cannot observe the movement of individual users between stages or infer any transition, because the two exports are aggregate, cross-sectional and very small.

Reading the Rankings as Multi-Stage Behavioural Performance

The psychologically informed lens defined in Section 2.6 determines how the rankings should be read. A high closeness coefficient identifies an alternative whose profile is comparatively balanced across observable funnel criteria at increasing levels of behavioural commitment: the leading organic post is neither the most widely exposed nor the most intensely liked, but the one that combines acceptable exposure efficiency with the strongest movement into purposive action, and the leading campaign is efficient at every observed depth from clicks through messages to follows. The coefficient is therefore a statement about the shape of a behavioural profile, not about a favourable psychological state; it does not indicate that an audience was more attentive, more reassured or more willing to enroll.

This restraint is of practical value precisely in daycare marketing. Choosing childcare is a consequential, context-sensitive family decision in which caregivers weigh service quality alongside promotional information (Albastho et al., 2024), so an account owner who reads a high engagement

rate as evidence of parental trust risks investing in content that generates cheap reactions rather than enquiries. Ordering the criteria by behavioural commitment makes the scarce step visible instead: in this account, acquisition rather than exposure limits the funnel. No caregiver-level construct is observed in these exports and none is claimed; the interpretive vocabulary stays at the level of the behavior the platform records. Systematic coding of message appeals, for example as reassurance about safety, demonstration of learning activity, community belonging or administrative announcement, would let content attributes enter the matrix explicitly, but the present reading does not depend on it.

Threats to Validity

The following threats are material and none of them is fully mitigated. Statistical conclusion validity is limited by sample size: ten organic alternatives and three campaigns support no inferential test, and CRITIC correlations computed on ten observations are themselves unstable (Schönbrodt and Perugini, 2013). Construct validity is limited because platform metrics are estimated and partly in development (Meta Platforms, n.d.-a, n.d.-b), and because one reach value is missing rather than zero. Internal validity is absent by design: the study is observational and posting time, content type, budget, duration, creative material and targeting vary together, with campaign periods overlapping in late November 2025. Temporal coverage is imbalanced, with five Wednesday posts, a single Saturday post and no Sunday post in the original export, and the time zone of the exported timestamps is not documented, so hour-level statements are unreliable. The advertising campaigns were all published in the afternoon or evening, so no time-of-day contrast exists. Finally, the outcome that matters commercially, namely profile visits, enquiries and enrolment, is not observed at all: follows and messaging conversations are proxies for it.

Generalisation

No result reported here generalises beyond this account and this window: the weights are functions of this decision matrix, the rankings are relative to the alternatives present, and the dominance frequencies describe the weight space of these criteria sets. What is transferable is the pipeline, whose funnel definitions, two-track weighting rule, ranking procedure and mandatory sensitivity report can be applied to any account whose owner can export the same platform-native counts.

Limitations and Future Work

The limitations of this proof of concept are stated deliberately and without mitigation claims:

1. Sample size. Ten complete-case organic posts and three campaigns permit description and ranking only; no inferential or predictive claim is possible.
2. Single account. One Indonesian daycare account in one city provides no basis for comparison across providers, market segments or countries.
3. Missing exposure value. One post report zero accounts reached with positive views and is excluded from the reach-based matrix; the exclusion itself is a judgement, and the alternative handlings change the CRITIC weights.
4. Imbalanced calendar coverage. Five of the eleven original posts fall on Wednesday, only one on Saturday and none on Sunday, so weekly comparison is impossible.
5. Undocumented time zone. The exported timestamps carry no time zone, so hour-level and time-band interpretations are unreliable.
6. Estimated platform metrics. Views, reach and interaction aggregates are documented by the platform as estimated and partly in development, which propagates into every ratio.

7. Unstable weighting inputs. CRITIC correlations at ten observations are far from the sample sizes at which correlation estimates stabilise, so the weights should be regarded as provisional.
8. Equal-weight benchmark for campaigns. The advertising weights are a neutral convention, not a measured preference, and a different stakeholder preference could be defended.
9. Uncontrolled campaign comparison. Budget, duration, creative material, targeting and calendar period vary jointly, and two campaigns overlap in time.
10. Boundary of the psychologically informed reading. Behavioural theory organises and disciplines the interpretation of the criteria, but no psychological construct was directly measured and no content coding was performed, so any statement about caregiver states lies outside the evidence base of this study.
11. No observable funnel transitions. The exports contain no user-level journey identifier, no repeated history for the same user, no purchase or enrolment outcome and no temporal sequence of stage changes, so no claim is made about detecting or predicting funnel transitions in the sense of Crowe et al. (2025).
12. No commercial outcome. Profile visits, enquiries and enrolments were not available, so funnel completion is proxied by follows and messaging conversations.

Future work follows directly from these limitations. The immediate priority is a longitudinal collection design covering at least one full year with balanced representation of every day of the week and every time band, recorded with an explicit time zone, and ideally spanning several comparable accounts. On that basis the pipeline can be extended in three directions: adding coded content attributes as criteria, so that message appeal enters the decision matrix explicitly; complementing objective CRITIC weights with elicited owner preferences in a hybrid index, as suggested by Diakoulaki et al. (1995); and linking the funnel to commercial outcomes by recording profile visits, enquiries and enrolments, which would allow the ranking to be validated against a decision-relevant criterion. A registered comparison of campaigns that differ in one factor at a time would be required before any scheduling or targeting recommendation could be made. Richer data collection, including coded message appeals or caregiver-level constructs measured with validated instruments, would widen the interpretation, although the behavioural reading adopted here does not require it.

CONCLUSION

This article presented an Instagram marketing analytics pipeline for micro-scale accounts that combines funnel-based criteria with objective weighting, distance-to-ideal ranking and exhaustive weight sensitivity analysis, and demonstrated it on one Indonesian daycare account. For the ten complete-case organic posts, CRITIC assigned weights of 0.2764 to reach efficiency, 0.3887 to engagement rate and 0.3349 to follow conversion, and TOPSIS placed the 15 May 04:02 post first with a closeness coefficient of 0.7503; the enumeration of all 231 admissible weight vectors qualified that result immediately, since the post leads in only 59.74 per cent of them. For the three completed campaigns, an equal-weight benchmark over four objective-aligned criteria placed the 22 to 27 November 2025 campaign first with a closeness coefficient of 1.0000 and a dominance frequency of 100 per cent, which reflects Pareto dominance on those criteria rather than causal evidence about scheduling. No optimal posting time, weekly pattern or causal mechanism is claimed. Psychological and consumer-behavior theory informed the criterion architecture and the vocabulary of interpretation, ordering the criteria by increasing behavioural commitment, while CRITIC, TOPSIS and FRWSA operated solely on observed platform metrics; no psychological state was measured or inferred. The

contribution is methodological: a transparent and reproducible way for very small accounts to reach defensible content and campaign decisions while stating precisely how much confidence their data can support.

REFERENCES

- Albastho, M. T., Nurtjahjani, F., Niaga, A., & Malang, P. N. (2024). Pengaruh promosi media Instagram dan kualitas pelayanan terhadap keputusan pembelian jasa penitipan anak di Abthre Daycare Kota Malang. *Jurnal Aplikasi Bisnis*, 10(2), 0–5. <https://doi.org/10.33795/jab.v10i2.1386>
- Artura, I. P., Anreano, F., Putri, H., Amanatul, I., Adiluhung, K., Ulya, C., & Gumelar, N. A. (2024). Pengaruh strategi marketing aplikasi e-commerce terhadap perilaku konsumtif mahasiswa Teknik Industri UNS dengan pendekatan psikologi konsumen. *Jurnal STIA Bengkulu: Committee to Administration for Education Quality*, 10(1), 33–40. <https://doi.org/10.56135/jsb.v10i1.133>
- Betri, T. J., Fadilah, H. N., & Nugroho, G. F. (2024). Data science analysis on UIN Raden Mas Said's media accounts, 8(1), 76–85. <https://doi.org/10.20961/ijscs.v7i2.96733> (journal title not stated in the copy consulted, and the volume numbering differs from the digital object identifier; authors should verify against the publisher record before final submission)
- Crowe, C., Haas, C., & Hall, M. (2025). A novel data-driven approach to detect and predict customer transitions in the marketing funnel. *Social Network Analysis and Mining*, 15, Article 105. <https://doi.org/10.1007/s13278-025-01533-9>
- Dewi, N. F., Riyadi, D., & Kusumo, R. (2023). Social media as a platform for information: Study in the marketing department at Hospital X. *Proceedings of the International Conference on Vocational Education and Applied Science and Technology (ICVEAST)*. <https://doi.org/10.2991/978-2-38476-132-6>
- Diakoulaki, D., Mavrotas, G., & Papayannakis, L. (1995). Determining objective weights in multiple criteria problems: The CRITIC method. *Computers & Operations Research*, 22(7), 763–770. [https://doi.org/10.1016/0305-0548\(94\)00059-H](https://doi.org/10.1016/0305-0548(94)00059-H)
- Fischer, G. W. (1995). Range sensitivity of attribute weights in multiattribute value models. *Organizational Behavior and Human Decision Processes*. <https://scholars.duke.edu/publication/771230> (volume and page range not stated on the record consulted)
- Gaona, Guisasaola, and Baeza (2025). An integrated TOPSIS framework with full-range weight sensitivity analysis for robust decision analysis. *Decision Analytics Journal*, 17, 100642. <https://ddd.uab.cat/pub/artpub/2025/330424/1-s2.0-S2772662225000980-main.pdf>
- Golder, S. A., & Macy, M. W. (2011). Diurnal and seasonal mood vary with work, sleep, and daylength across diverse cultures. *Science*, 333(6051), 1878–1881. <https://doi.org/10.1126/science.1202775>
- Gräve (2019). What KPIs are key? Evaluating performance metrics for social media influencers. *Social Media + Society*. <https://journals.sagepub.com/doi/10.1177/2056305119865475>
- Hibel, L. C., Mercado, E., & Trumbell, J. M. (2012). Parenting stressors and morning cortisol in a sample of working mothers. *Journal of Family Psychology*, 26(5), 738–746. <https://doi.org/10.1037/a0029340>
- Kaur, P., Dhir, A., Chen, S., & Rajala, R. (2016). Flow in context: Development and validation of the flow experience instrument for social networking. *Computers in Human Behavior*, 59, 358–367. <https://doi.org/10.1016/j.chb.2016.02.039>

- Mahatmi, M. W., Iswanti, S., & Hanafi, M. N. (2022). Pemanfaatan media sosial Instagram sebagai sarana promosi di Twinkle Daycare & Courses. *Jurnal Pengabdian Masyarakat – Teknologi Digital Indonesia*, 1(2), 57. <https://doi.org/10.26798/jpm.v1i2.596>
- Pencarelli, T., & Mele, M. G. (2019). A systematic literature review on social media metrics. *Mercati & Competitività*, 2019(1), 15–38. <https://doi.org/10.3280/mc1-2019oa7624>
- Qin, V., Pauwels, K., & Zhou, B. (2024). Data-driven budget allocation of retail media by ad product, funnel metric, and brand size. *Journal of Marketing Analytics*, 12, 235–249. <https://doi.org/10.1057/s41270-024-00294-2>
- Şahin (2020). A comprehensive analysis of weighting and multicriteria methods in the context of sustainable energy. *International Journal of Environmental Science and Technology*, 18(6), 1591–1616. <https://pmc.ncbi.nlm.nih.gov/articles/PMC7490576/>
- Sano, Y., Takayasu, H., Havlin, S., & Takayasu, M. (2019). Identifying long-term periodic cycles and memories of collective emotion in online social media (manuscript, pages 1–17; publication venue not stated in the copy consulted).
- Schönbrodt, F. D., & Perugini, M. (2013). At what sample size do correlations stabilize? *Journal of Research in Personality*. <https://www.psy.lmu.de/gp/news/corevol2/index.html> (volume and page range not stated on the record consulted)
- Töllinen, A., & Karjaluo, H. (2011). Marketing communication metrics for social media. *International Journal of Technology Marketing*, 6(4), 316–330. <https://scispace.com/pdf/marketing-communication-metrics-for-social-media-17up65e1vs.pdf>
- Trunfio, M., & Rossi, S. (2021). Conceptualising and measuring social media engagement: A systematic literature review. *Italian Journal of Marketing*, 2021(3), 267–292. <https://pmc.ncbi.nlm.nih.gov/articles/PMC8354841/>
- von Nitzsch, R., & Weber, M. (1993). The effect of attribute ranges on weights in multiattribute utility measurements. *Management Science*. <https://pubsonline.informs.org/doi/10.1287/mnsc.39.8.937> (publication year not stated on the record consulted)
- Wedley et al. (1999). Interpretation of criteria weights in multicriteria decision making. *Computers & Industrial Engineering*. <https://www.sciencedirect.com/science/article/abs/pii/S036083520000019X> (volume and page range not stated on the record consulted)
- Wegner et al. (2023). Performance analysis of social media platforms: Evidence of digital marketing. *Journal of Marketing Analytics*. <https://pmc.ncbi.nlm.nih.gov/articles/PMC9936493/> (complete author list, volume and page range not stated in the copy consulted; authors should verify against the publisher record before final submission)
- Zhu, Y., Tian, D., & Yan, F. (2020). Effectiveness of entropy weight method in decision-making. *Mathematical Problems in Engineering*, 2020, 3564835. <https://doi.org/10.1155/2020/3564835>