

What Mortality-Reminder Experiments Identify About Motives: Partial Identification and the Economic Value of Motive Evidence

Hamed Mohammadpour ¹, Ghahreman Abdoli ^{2,*}



¹ Faculty of Economics, University of Tehran, Tehran, Iran; 

² Faculty of Economics, University of Tehran, Tehran, Iran; 

* Correspondence: abdoli@ut.ac.ir

Citation: Mohammadpour, H., & Abdoli, G. (2027). What Mortality-Reminder Experiments Identify About Motives: Partial Identification and the Economic Value of Motive Evidence. *Business, Marketing, and Finance Open*, 4(7), 1-21.

Received: 05 August 2026

Revised: 15 August 2026

Accepted: 23 August 2026

Initial Publication: 26 August 2026

Final Publication: 01 December 2027



Copyright: © 2027 by the authors. Published under the terms and conditions of Creative Commons Attribution-NonCommercial 4.0 International (CC BY-NC 4.0) License.

Abstract: Under standard design conditions, randomization identifies a mortality-reminder assignment effect on a prespecified economic choice, not why choice changed. The analysis separates this effect from arm-marginal changes in theory-indexed practical content and individual component transitions. Inferring noninstrumental concern for a specified other's good or the actor's own good requires defensible cross-arm alignment and four bridges: content, choice, measurement, and causal order. Under the paired-construction closure condition, other-good and self-good specifications can generate the same arm-specific observable-law family. Valid aligned component marginals yield sharp marginal Fréchet bounds on individual transitions, but not the joint transition law. Narrowing the supported state set may inform without changing action; motive evidence has economic decision value only if it changes a prespecified prediction, transport conclusion, policy ranking, minimax-regret choice, or welfare conclusion. A consequential-giving construction shows states sharing one baseline choice law can reverse the ranking of beneficiary-impact versus public-recognition fundraising policies. The dated, search-bounded map contains 221 records from 121 works. Of 93 source-version records in a post hoc five-gate audit, 18 records from 13 works met all criteria. Because the audited studies lack a common estimand, results are not pooled or vote-counted; the audited evidence supports no domain-general shift toward altruism or self-regard. The proposed eight-cell giving design varies cue, beneficiary impact, and truthful publicity; an actor-return-substitute follow-up tests prespecified causal-order predictions after independent validation that the substitute generates the targeted actor return. Under valid exclusions and decision-calibrated precision, the design can discriminate among economic-component restrictions; neither design point-identifies unrestricted practical motives.

Keywords: motive identification; mortality reminders; partial identification; economic choice; charitable giving

1. Introduction

Mortality reminders occupy an unusual position at the intersection of psychology, behavioral economics, and decision science because they are experimentally tractable while potentially affecting a wide range of economically meaningful choices. The mortality-salience literature, developed largely within terror management research, has examined whether making death cognitively accessible changes evaluations, social judgments, consumption, intergroup responses, risk preferences, temporal choices, and prosocial behavior. Early quantitative synthesis suggested that mortality-salience manipulations could generate systematic psychological and behavioral effects across diverse domains [1]. More recent reassessments, however, have placed substantially

greater emphasis on publication bias, replication, study heterogeneity, and the specificity of experimental contexts, weakening the basis for treating “mortality salience” as producing a single stable behavioral response [2]. In consumer research specifically, systematic review and meta-analytic evidence likewise indicates a heterogeneous literature in which outcomes depend on the consumer domain, manipulation, theoretical moderator, and behavioral measure [3]. This heterogeneity creates an important distinction between establishing that an experimentally assigned mortality reminder changes a particular choice and establishing what psychological or economic motive produced that change.

The empirical literature illustrates the breadth of the problem. Mortality reminders have been associated with charitable behavior, prosocial intentions, fairness-related decisions, saving, retirement choices, delayed rewards, status consumption, and intergenerational allocations. Mortality salience has altered charitable behavior in ways contingent on social and normative context [4], while mortality-related manipulations have also been linked to blood-donation intentions and other prosocial responses [5, 6]. Evidence from a national sample further suggests that mortality salience can affect charitable allocations, although an observed increase in giving remains a behavioral outcome rather than direct evidence of an altruistic motive [7]. More recent economic experimentation has similarly examined whether mortality salience changes prosocial behavior under controlled incentives [8], whereas research distinguishing death reflection from death anxiety suggests that different forms of mortality-related cognition may produce divergent routes to prosociality [9]. Pandemic-related mortality cues add further contextual complexity because they combine mortality awareness with threat, collective identity, uncertainty, and social connection or isolation [10]. Endpoint-focused meditation also illustrates how contemplating finitude can be embedded within a broader cognitive and emotional intervention rather than functioning as a simple, isolated death cue [11]. Consequently, the experimental label attached to a mortality manipulation cannot by itself determine the content or mechanism through which behavior changes.

The same diversity appears in economic choices that do not involve an obvious beneficiary. Mortality salience has been investigated in saving behavior and financial decision-making, with evidence that financial accumulation itself may acquire a psychologically protective function under mortality threat [12]. Retirement-decumulation research has shown that mortality cues and the wording of annuity options can influence preferences over lifetime income streams [13], while experimental work on commercial annuities and bequests further demonstrates that mortality-related framing may affect choices involving consumption and legacy provision [14]. Intergenerational research suggests that mortality awareness can increase attention to legacy and future others, thereby affecting decisions whose consequences extend beyond the actor's lifetime [15]. Yet replication evidence in risk and intertemporal choice has not consistently confirmed previously proposed mortality-related effects [16, 17]. Likewise, recent preregistered work on status consumption provides a stringent test of mortality-salience predictions without establishing a universal mortality-induced shift toward status seeking [18]. These findings caution against moving directly from a treatment effect on one observed choice to a general conclusion about self-regard, altruism, materialism, impatience, or legacy motivation.

This inferential problem is especially acute in prosocial settings because behavior and motive are conceptually distinct. Prosocial behavior refers to actions that benefit others, whereas altruism concerns the motivational basis for such action and therefore cannot be inferred merely from the distribution of material consequences [19]. Philosophical analyses have similarly emphasized that altruism is not a simple behavioral label but involves substantive questions about whose good functions as an end and how competing self- and other-directed reasons enter practical deliberation [20]. Indeed, altruism has been characterized as a “thick” concept whose application

combines descriptive and evaluative elements, making its operationalization more demanding than classifying a transfer as generous or prosocial [21]. Contemporary work on human motivation further underscores that pleasure, affect, concern for others, and self-directed benefits can coexist, and that the presence of an enjoyable or relieving consequence does not automatically establish hedonistic motivation [22]. Thus, a mortality reminder may increase a donation while leaving open whether the actor is guided by the beneficiary's good, the value of performing the act, anticipated emotional relief, moral duty, identity maintenance, social approval, or some combination of these considerations.

Economic theories of giving provide a disciplined vocabulary for distinguishing some of these possibilities, but they do not automatically settle the deeper motivational interpretation. Models of pure altruism, warm glow, reciprocity, fairness, social norms, and mixed motives show that observationally similar transfers may arise from structurally different preferences [23]. Experimental variation in outside provision has been used to distinguish pure from impure altruism and to identify the relative importance of beneficiary output and the utility associated with giving itself [24]. Related work has sought to bridge economic and psychological approaches to altruism and egoism by clarifying how apparently similar constructs are operationalized across traditions [25]. More recent synthesis emphasizes the need to distinguish economic components of giving from broader psychological interpretations of why those components matter to an actor [26]. These distinctions are crucial for mortality-reminder experiments because an increase in the marginal value of giving may reflect concern for recipient outcomes, but it may also reflect identity, agency, duty, anticipated relief, or another actor-side return.

Warm-glow explanations illustrate this ambiguity particularly clearly. Warm glow is frequently invoked when individuals value the act of giving independently of its incremental effect on recipient welfare, but this behavioral component does not uniquely identify a self-regarding practical end. Conceptual analysis has questioned whether experiencing warm glow after giving actually explains why the person gave or merely describes an affective correlate or consequence of the action [27]. Experimental evidence on the phenomenology of warm glow likewise shows that positive affect associated with giving requires careful interpretation with respect to its timing, determinants, and causal role [28]. Behavioral validation of warm-glow questionnaires improves measurement of the construct but does not eliminate the need to specify how a measured construct maps onto the motive under investigation [29]. Likewise, experimental measures of the strength of altruistic motives can sharpen comparisons across individuals and tasks while remaining conditional on the behavioral and theoretical structure that gives those measures meaning [30]. Measurement validity is therefore necessary but not sufficient for identifying the practical content that guided a choice.

Other experimental margins demonstrate how much can be learned when choice environments are designed to discriminate among competing explanations. Field experiments that vary solicitation conditions, advance warning, and opportunities to avoid a fundraiser distinguish willingness to give from social pressure and sorting into or away from an interaction [31]. More recent work on moral-preference elicitation shows how image concerns can alter observed prosocial choices and how experimental design can be structured to separate preference expression from reputational incentives [32]. Intertemporal experimental designs can also distinguish utility associated with the donor's choice from utility tied to when consequences reach a beneficiary [33]. Such designs reveal economically interpretable components, but the translation from an identified component to a practical motive still requires additional assumptions. For example, sensitivity to beneficiary impact may reflect concern for the beneficiary as an end, but it may also increase anticipated pride or meaning; similarly, observability may activate reputational utility,

accountability, conformity, or norm compliance. Preference-based interpretations can therefore become misleading when analysts treat observed choices as transparent expressions of latent welfare-relevant preferences [34].

The problem becomes still more important when mortality reminders alter the way options are construed. Reason-based choice theory permits the properties and reasons that are salient to a decision maker to depend on context rather than assuming that treatment simply changes fixed preference weights [35]. Under mortality salience, the same transfer could therefore be represented as helping another person, leaving a legacy, affirming a moral identity, reducing existential discomfort, complying with a norm, or restoring control. Evidence that mortality salience can increase satisfaction following prosocial behavior illustrates precisely why causal ordering matters: increased satisfaction may be an anticipated goal, a constitutive aspect of valuing the beneficiary's good, or a post-choice consequence rather than the motive that generated the choice [36]. Similarly, self-control has been shown to moderate mortality-related fairness decisions, highlighting how individual regulatory capacities can change observed treatment responses without uniquely revealing their motivational content [37]. These context-dependent pathways make a simple dichotomy between “altruistic” and “egoistic” responses empirically inadequate.

Replication and evidential diversity further complicate inference. Some preregistered or high-powered studies have failed to reproduce mortality-salience effects on established and novel behavioral measures [38], while other studies report effects that depend on the exact outcome, comparator, population, or manipulation. Such variation need not imply that all findings are uninformative; rather, it requires separating the estimand identified by each design from broader claims that the design does not warrant. Triangulation can strengthen inference when different methods have genuinely different vulnerabilities and nevertheless target the same phenomenon, but the mere accumulation of heterogeneous measures does not automatically identify a common latent construct [39]. A mortality-reminder literature spanning real donations, hypothetical intentions, risk tasks, annuity choices, status consumption, and affective reports therefore cannot be treated as though every study were a noisy measure of one underlying prosocial or self-regarding tendency.

These issues connect mortality-reminder research to the broader econometric problem of partial identification. In many empirical settings, available observations and maintained assumptions determine a set of admissible parameter values rather than a unique point, and scientifically meaningful analysis should characterize that identified set instead of imposing additional assumptions merely to obtain a scalar estimate [40]. Bayesian and frequentist approaches to partially identified models likewise emphasize the distinction between information supplied by the likelihood and conclusions supplied by priors or additional restrictions [41]. This logic is particularly appropriate for motive inference: randomized assignment can identify the causal effect of a mortality cue on an observed choice under standard experimental assumptions, yet several motive structures may remain observationally equivalent. The existence of a precise treatment effect therefore need not imply precise knowledge of the motive that generated it.

Recent advances in causal inference sharpen this distinction for latent outcomes and mechanisms. Methods for causal inference with latent outcomes formalize the assumptions required to infer treatment effects when the construct of interest is not directly observed [42], while nonparametric work on latent outcomes emphasizes measurement structures and bridge conditions needed to connect observed indicators to causal targets [43]. Mechanism analysis similarly cannot be reduced to documenting an intermediate correlation. Formal tests can assess hypotheses such as full mediation or bound effects operating outside a proposed mediator without thereby proving that the surviving pathway has a unique psychological interpretation [44]. Related work examines both

full mediation and the conditions under which causal mechanisms themselves are identifiable, reinforcing the principle that rejecting one mechanism is not equivalent to identifying another [45]. For mortality-reminder experiments, this implies that reports of affect, meaning, identity, death-thought accessibility, or anticipated satisfaction are informative only under explicit assumptions linking those measurements to the practical motive under dispute.

The distinction between behavioral identification and motive identification also matters because motive information is valuable only relative to a decision problem. Behavioral welfare economics has long recognized that observed choice need not provide an unambiguous welfare criterion when behavior is context-sensitive or influenced by ancillary psychological forces [46]. Partial identification extends this point by making uncertainty explicit rather than hiding it within a single preferred specification. Contemporary decision theory shows that treatment or policy choice can be formulated when payoff-relevant states are only partially identified, with optimal actions depending on the identified set and the chosen decision criterion [47]. Related work demonstrates how robust decision rules can be constructed when payoffs themselves are partially identified [48]. Consequently, learning more about motives has economic decision value not simply because the identified set becomes narrower, but because the additional information may change a prediction, a transport conclusion, a robust policy ranking, a minimax-regret decision, or a welfare judgment.

This decision-centered perspective is especially relevant to fundraising and other consumer-facing interventions. If mortality-linked giving primarily responds to verified beneficiary impact, improving the effectiveness of a donation may dominate increased publicity; if the response primarily reflects recognition or social image, public acknowledgment may instead be the operative margin; and if giving is substantially generated by pressure or avoidance costs, donor welfare and the availability of a private refusal option become essential considerations. Existing evidence is not sufficient to assume one of these interpretations across settings. The empirical record encompasses mortality-related changes in generosity, intentions, legacy decisions, satisfaction after giving, financial choices, and consumer responses, alongside null or contrary findings [4, 15, 18, 36, 38]. Accordingly, a framework is needed that preserves the causal information supplied by randomized mortality reminders while explicitly representing what remains unidentified about practical motivation.

The aim of this study was to determine what mortality-reminder experiments identify about economic motives by separating randomized choice effects from theory-indexed motive-component effects and individual motive transitions, and to establish when narrowing the resulting identified set changes economically relevant predictions, policy rankings, or welfare conclusions.

2. Methodology

The study employed a methodological and formal-analytic design combining partial-identification analysis, theory-indexed causal interpretation, a search-bounded evidence map, and a source-sensitive audit of mortality-reminder experiments. It was not a newly conducted behavioral experiment, clinical study, or participant-level secondary-data analysis. No new participants were recruited or analyzed, and the empirical component relied primarily on published or author-supplied study-level information and bibliographic metadata; publicly available source data were consulted only where necessary for source verification rather than for new participant-level inferential analyses. The unit of analysis therefore differed according to the analytical stage. For the formal identification analysis, the units were theoretical specifications, observable-law families, arm-specific economic choices, theory-indexed motive components, and admissible sets of data-generating processes. For the empirical

evidence map, the units were bibliographic manifestations, intellectual works, experiment/analysis records, and source-version records, which were deliberately kept distinct during reconciliation. The search was frozen on August 11, 2026. Retrieval through the official Crossref, OpenAlex, and PubMed interfaces, an OSF screening release, and a 67-record legacy atlas from an earlier phase of the project generated 6,013 route-hit incidences across 156 logged search routes. DOI-first clustering was used wherever DOI information was available, while exact title, publication year, and author matching was applied to records lacking a DOI. This process yielded 3,222 candidate manifestations, after which manual reconciliation produced a canonical evidence map containing 221 records representing 121 distinct works. Of these, 74 records from 40 works were verified from primary or accepted full texts, while 35 records from 29 works satisfied the stricter display rule used in the manuscript.

A frozen, source-reconciled subset was subsequently subjected to a post hoc five-gate audit designed as a methodological stress test rather than as a conventional meta-analytic inclusion procedure. Admission was determined independently of the direction, magnitude, or statistical significance of reported effects. The treatment criterion required randomized assignment to an explicit death-specific or mortality-risk cue against a nondeath comparator; broad compound shocks and designs containing inseparable co-manipulations were not treated as equivalent mortality-reminder contrasts. The outcome criterion required the absence of payoff-relevant partner response and classified eligible outcomes into actor-recipient or cause allocations and transfers, observed helping-directed behavior or stated prospective helping/willingness, or private risk, time, retirement, and consumption choices. Attitudinal outcomes, rankings without an act or willingness measure, fixed-wage performance, and payoff-interdependent strategic games were outside the focal audit. The source criterion required examination of the exact version of record or an accepted full-text version at the study level. The information criterion required either consequential or lottery-implemented choice, a randomized discriminatory margin such as price, cost, beneficiary output or need, recipient identity or horizon, norm, observability, avoidance, or death wording, or a time-stamped preregistered replication or confirmatory test with an a priori power or precision target. These information paths were distinguished as consequential implementation, discriminatory randomized variation, or preregistered independent precision/replication. Finally, the version criterion required reconciliation of record and work identities across multiple manifestations. The audit covered 93 study-level source/version records; after consolidation of multiple versions, 18 records representing 13 distinct works passed all criteria, whereas 75 failed at least one ordered gate. The qualifying records comprised 10 Tier-A studies from eight works and eight Tier-B studies from five works. Passing these criteria established eligibility only and was not interpreted as equal evidentiary weight. Because of project-imposed retrieval limits, machine-assisted preliminary exclusions, unresolved source access, and heterogeneous study designs, the evidence map was explicitly treated as search-bounded rather than as a census or conventional systematic review, and the study did not claim exhaustive coverage of the mortality-reminder literature.

Data collection was undertaken through a combination of bibliographic retrieval systems, source-verification procedures, structured coding rules, and formal theory-indexed classification tools rather than psychometric questionnaires or newly administered experimental instruments. The bibliographic component used the official Crossref, OpenAlex, and PubMed interfaces, supplemented by an OSF screening release and the pre-existing 67-record legacy atlas. The supplementary audit infrastructure recorded the exact search queries, encoded search URLs, date restrictions, retrieval caps, numbers of records returned, and execution status for each search route, thereby providing a version-controlled account of how the search-bounded evidence map was assembled. The workflow combined API-based retrieval, parser audits, a versioned bibliographic ledger, manual record

reconciliation, and scripted checks of denominators and algebraic quantities. Candidate records, intellectual works, experimental or analytical records, and effect-size incidences were retained as distinct entities to avoid conflating multiple publications or manifestations of the same underlying study. Central numerical claims were retained only after validation against the strongest available source basis, and source-access fields distinguished primary or accepted-full-text verification from secondary, abstract-based, or repository-based extraction. The machine-readable audit records preserved variables relevant to methodological interpretation, including treatment, comparator, temporal delay, recipient reality, implementation status, source version, and the first audit gate failed by an excluded record. This source-sensitive approach was particularly important because the manuscript did not assume that all mortality-reminder studies estimated a common construct or common causal estimand.

The conceptual measurement framework was based on three distinct vocabularies: material consequences, economic-choice components, and practical content. Material coding described who received an economic consequence; economic-choice coding captured components such as public-good provision, own-gift value, impact sensitivity, observability, avoidance, and timing; and practical-content coding concerned what had operative noninstrumental standing in guiding the actor. The two focal theory-indexed components were other-good content, denoted O , and self-good content, denoted S . O was coded when the good of a specified other person or determinate aggregate had an operative noninstrumental role in guiding behavior, whereas S applied when the actor's own good occupied that role. The coding system also retained principle-related content, other social content, and a residual category, with the first four categories permitted to overlap. The record further distinguished locus or content type, value status such as noninstrumental, instrumental, enabling, incidental, or inert status, and causal role such as independently sufficient, jointly necessary, operative-but-neither, or inert. Because a mortality cue could alter the meaning or construal of an option rather than merely change the weight assigned to an invariant motive, cross-arm comparisons required explicit alignment of the content bearer, represented good, temporal horizon, level of granularity, institutional context, and criterion of practical control.

Four bridge conditions were treated as the central inferential instruments linking randomized assignment and observable behavior to motive-relevant conclusions. The content bridge specified the theory of practical content and the validity of cross-arm alignment. The choice bridge specified how contents, beliefs, constraints, attention, and stochastic choice processes generated observed action. The measurement bridge linked self-reports, affective responses, and process measures to the target construct while explicitly allowing reactivity and differential measurement error. The causal-order bridge distinguished an actor-side benefit that was prospectively sought as an end from satisfaction that was constitutive, jointly ultimate, a downstream by-product, a reinforcing consequence, or an effect of a common cause. The manuscript emphasized that temporal precedence alone did not establish instrumental or motivational order. Operational validation therefore considered allocation-procedure fidelity, attrition, baseline-balance diagnostics, independent construal coding, comprehension, incentive consistency, validation samples, temporal priority, method-divergent convergence, and, where applicable, randomized actor-return substitutes or agency manipulations. These procedures did not automatically establish O or S ; rather, they progressively restricted the set of practical-content structures compatible with the observed evidence.

Data analysis combined causal estimand definition, partial identification, observational-equivalence analysis, sharp bounding, decision analysis, and a nonpooled qualitative audit of the empirical literature. The first analytical target was the randomized intention-to-treat effect of mortality-reminder assignment on a prespecified economic choice. Let $Z_i \in \{0,1\}$ denote randomized assignment to the comparator and mortality-reminder conditions,

respectively, and let $Y_i(z)$ denote the corresponding potential economic choice. The choice-effect estimand was defined as

$$\tau_Y = E[Y_i(1) - Y_i(0)].$$

Under random assignment, consistency, positivity, absence of relevant interference, valid construction of the observed outcome, and complete observation—or an explicitly specified attrition estimand with appropriate identifying or bounding assumptions—this quantity can be identified or bounded for the randomized population. The analysis deliberately separated this behavioral effect from any conclusion about why the behavior occurred, because randomization identifies the contrast in choice but does not by itself establish the truth of the content, measurement, choice, or causal-order bridges.

For motive-relevant analysis, arm-specific practical-content variables were aligned to a common theory-indexed content space. The aligned record can be represented generically as

$$M_i(z; \pi, \omega) = [O_i(z; \pi, \omega), S_i(z; \pi, \omega), P_i(z; \pi, \omega), C_i^{\text{soc}}(z; \pi, \omega), X_i(z; \pi, \omega)] \in \mathcal{M}^* \subseteq \{0,1\}^5.$$

Here, O denotes other-good content, S self-good content, P principle content, C_i^{soc} other social content, and X the residual category; π represents the cross-arm alignment and ω an admissible structural specification. The first four coordinates are nonexclusive, while the residual coordinate is defined by the absence of all four coded classes. A typical arm-marginal component effect for other-good content was

$$\tau_O(\pi, \omega) = P\{O_i(1; \pi, \omega) = 1\} - P\{O_i(0; \pi, \omega) = 1\}.$$

The individual-level transition target was the joint distribution

$$Q_M^{\pi, \omega} = \mathcal{L}\{M_i(0; \pi, \omega), M_i(1; \pi, \omega)\}.$$

This joint transition law cannot generally be recovered from separate treatment-arm marginals because those marginals contain no information about the within-person coupling of potential motive states. Consequently, alignment was treated as logically prior to statistical identification: if no defensible cross-arm map existed, the relevant scalar motive-effect estimand was considered undefined rather than merely imprecisely estimated.

Formal identification was evaluated using observable-law fibers and identified sets. Let $\mathcal{A}$ denote the maintained structural assumptions, Π the admissible alignment class, T the theory or bridge restrictions, and P^* the observed population law. The observational fiber was defined as

$$F(P^*; \mathcal{A}, \Pi, T) = \{\omega \in \Omega(\mathcal{A}, \Pi, T) : P_\omega = P^*\}.$$

For a target functional $t(\omega)$, the corresponding identified set was

$$\Theta_t(P^*; \mathcal{A}, \Pi, T) = \{t(\omega) : \omega \in F(P^*; \mathcal{A}, \Pi, T)\}.$$

A target was regarded as point identified only when $t(\omega)$ was constant over a nonempty observational fiber. This formulation was used to distinguish statistical precision from structural identification: greater precision in estimating the observable distribution cannot, by itself, shrink an identified set generated by unresolved observational equivalence. The central nonidentification analysis therefore examined paired O - and S -based specifications that preserved the same observable kernels and joint observable distribution while assigning different practical-content interpretations to the underlying choice process. Under the stated paired-construction closure condition, differences in arm prevalence, component effects, or coordinate-specific transition laws were not point identified even though assignment, choice, beneficiary output, affect, impact-belief reports, and reason reports could have identical observable distributions.

When stronger assumptions justified valid measurement of an aligned binary component in both arms, the study derived sharp marginal Fréchet bounds for individual transitions. Let

$$p_0 = P\{O_i(0) = 1\}, \quad p_1 = P\{O_i(1) = 1\}.$$

Let q_{01} and q_{10} denote the probabilities of movement from $O = 0$ to $O = 1$ and from $O = 1$ to $O = 0$, respectively. The sharp scalar bounds were

$$\begin{aligned} \max(0, p_1 - p_0) &\leq q_{01} \leq \min(1 - p_0, p_1), \\ \max(0, p_0 - p_1) &\leq q_{10} \leq \min(p_0, 1 - p_1), \end{aligned}$$

with the overall probability of an individual-level change bounded by

$$|p_1 - p_0| \leq P\{O_i(0) \neq O_i(1)\} \leq \min(p_0 + p_1, 2 - p_0 - p_1).$$

These bounds were interpreted as sharp separately for each scalar transition quantity, not as a Cartesian joint identified set. A single coupling table must satisfy all transition restrictions simultaneously, and coordinatewise bounds for O and S do not identify the complete joint transition law across the five-component content vector.

The economic relevance of motive information was analyzed by examining whether additional evidence changed a prespecified decision object rather than merely narrowing the set of admissible explanations. For an admissible state set Ω , a finite policy set D , and welfare function $W(a, \omega; v, h)$, the set of policies optimal in every admissible state was represented as

$$A^0(\Omega) = \bigcap_{\omega \in \Omega} \operatorname{argmax}_{a \in D} W(a, \omega; v, h).$$

When no policy was uniformly optimal, ambiguity was evaluated using minimax regret. For policy a ,

$$\operatorname{Reg}(a; \Omega) = \sup_{\omega \in \Omega} \left[\max_{a' \in D} W(a', \omega; v, h) - W(a, \omega; v, h) \right],$$

and the minimax-regret policy set was

$$A^{MR}(\Omega) = \operatorname{argmin}_{a \in D} \operatorname{Reg}(a; \Omega).$$

Thus, contraction of an identified or supported state set was considered economically valuable only if it changed a prediction, transport conclusion, policy ranking, welfare conclusion, robust-dominance relation, or regret-based decision. Mere reduction in interpretive uncertainty was not equated with decision value.

Finally, an illustrative consequential-giving model was used to show how behaviorally equivalent mortality-reminder effects could imply different policy recommendations. With G denoting a binary transfer, m mortality-reminder assignment, κ beneficiary output per transfer, v truthful external observability, and $\Lambda(\cdot)$ the standard logistic distribution function, the normalized choice model was

$$p_{m\kappa v} = P(G = 1 \mid m, \kappa, v) = \Lambda(-2 + a_m \kappa + b_m + s_m v),$$

where a_m represented beneficiary-output sensitivity, b_m impact- and publicity-invariant own-act value, and s_m observability sensitivity. The mortality-cue contrast on the normalized logit index was

$$D_{\kappa v} = \operatorname{logit}(p_{1\kappa v}) - \operatorname{logit}(p_{0\kappa v}) = \Delta a \kappa + \Delta b + \Delta s v.$$

This model was used as an existence construction rather than as an empirical calibration. It demonstrated that an identical observed two-arm baseline effect could be consistent with shifts in beneficiary-output sensitivity, own-act value, or visibility sensitivity and could therefore reverse the ranking of beneficiary-impact and public-recognition policies. The empirical evidence map was analyzed separately from these formal constructions. Because the admitted studies differed substantially in mortality cues, comparators, populations, timing, consequentiality, prices, observability, outcomes, and estimands, no pooled effect size, sign count, or domain-general average was calculated. Study-level results were instead reported with their design characteristics, implementation status, source-verification status, and maximum warranted inference, preserving the distinction between evidence of a choice effect and evidence concerning its underlying practical motive.

3. Findings and Results

Theorem 1 (nonidentification under incomplete bridge restrictions). Fix the experimental support $\mathcal{K}$ of recipient impact. For every admitted O specification whose disputed term satisfies $S_z^O = 0$ and has scalar choice index $d_z^O = b_z + O_z v_z$, suppose the admissible class is closed under the following paired construction:

1. set $O_z^S = 0$, $S_z = O_z$, and $d_z^S = b_z + S_z \ell_z$;
2. set $\ell_z = v_z$ on $\mathcal{K}$; and
3. copy all other content coordinates, the context and heterogeneity law, and the choice, downstream-consequence, belief, affect, and measurement kernels.

The paired O and S specifications then induce the same full arm-specific observable-law family. Any well-defined target that differs between the paired specifications—such as the relevant arm prevalence, arm effect, or coordinate-labeled component-transition law—is therefore not point identified in that class.

This is not a global impossibility result. Impact variation separates the rivals only if actor return cannot mimic impact value; a prospective report does so only if its measurement model bears on instrumental order beyond shared dependence on choice and affect. The closure condition is substantive: the paired specification must remain admissible under the same content, choice, measurement, and causal-order restrictions; it is not generated by relaxing those bridges after observing the data. The supplement provides the formal proof at the level of the observable kernels.

For a binary gift Y , fixed $\kappa_0 > 0$, $\delta > 0$, and standard-logistic ε_z , let

$$Y(z) = \mathbf{1}\{-c + z\delta\kappa_0 + \varepsilon_z \geq 0\},$$

$$H(z) = \kappa_0 Y(z),$$

$$A(z) = a_0 + \delta H(z) + \eta_z,$$

$$\widehat{\kappa}(z) = \kappa_0 + \nu_z,$$

$$J(z) = j\{Y(z), H(z), A(z), \widehat{\kappa}(z)\} + \xi_z.$$

Under controlled gift g , let $A_z(g) = a_0 + \delta\kappa_0 g + \eta_z$, pre-choice information $\mathcal{I}_z$, and forecast actor affect $L_z(g) = E\{A_z(g)|\mathcal{I}_z\} - a_0 = \delta\kappa_0 g$. In DGP O,

$$M^O(0) = (0, 0, 0, 0, 1), \quad M^O(1) = (1, 0, 0, 0, 0).$$

Recipient output has operative noninstrumental control, ordered treatment $\rightarrow$ O-content $\rightarrow Y \rightarrow H \rightarrow A$; L may be accurate but is not an end. Thus $(\tau_O, \tau_S) = (1, 0)$. In DGP S,

$$M^S(0) = (0, 0, 0, 0, 1), \quad M^S(1) = (0, 1, 0, 0, 0).$$

Anticipated actor relief has operative noninstrumental control, ordered treatment $\rightarrow$ S-content (prospective L) $\rightarrow Y \rightarrow H \rightarrow A$; recipient output instruments relief. Thus $(\tau_O, \tau_S) = (0, 1)$. Both DGPs impose $E(\eta_z|\mathcal{I}_z) = 0$, identical remaining disturbance dependence, observable mapping, and joint disturbance law. Assignment, gift, output, affect, impact-belief report, and reason report consequently have the same joint distribution although targets differ. Randomizing κ separates them only by excluding anticipated relief that rises with effective help; the exclusion, not randomization alone, supplies force.

Under the stronger assumption that aligned binary O-components are validly measured in both arms and $p_z = P\{O(z) = 1\}$ is identified. Let $q_{01} = P\{O(0) = 0, O(1) = 1\}$ and $q_{10} = P\{O(0) = 1, O(1) = 0\}$.

Proposition 2 (sharp Fréchet bounds on individual component transitions). This is a mortality-component specialization of classical Fréchet bounds for marginals-only treatment responses [37], not a new generic bounding result.

$$\begin{aligned} \max(0, p_1 - p_0) &\leq q_{01} \leq \min(1 - p_0, p_1), \\ \max(0, p_0 - p_1) &\leq q_{10} \leq \min(p_0, 1 - p_1), \\ \text{and} \\ |p_1 - p_0| &\leq P\{O(0) \neq O(1)\} \leq \min(p_0 + p_1, 2 - p_0 - p_1). \end{aligned}$$

Each displayed scalar bound is sharp in isolation: every value within each interval is attainable under some admissible coupling, but values across the different transition quantities are not freely jointly attainable. Thus $\tau_O = p_1 - p_0 = q_{01} - q_{10}$ is identified, not the number or direction of individual changes. The displayed intervals are not a Cartesian joint set: one 2×2 coupling table jointly constrains the transition quantities. Coordinatewise O/S application gives marginal bounds, not a coherent set for joint five-coordinate $\mathcal{Q}_M$, which needs higher-dimensional coupling respecting nonexclusivity and the residual rule. Monotone switching may tighten one coordinate. Repeated-measure crossover sharpens coupling only with period stability, no anticipation or carryover, and a map preserving bearer, represented good, horizon, institutional meaning, and decision context. The supplement constructs the bounds.

Let behavioral evidence leave Ω_B , motive-relevant evidence leave $\Omega_{B+M} \subseteq \Omega_B$, and $\Gamma_q(\Omega) = \{q(\omega): \omega \in \Omega\}$ for conclusion q . Strict contraction is informative but need not change action.

For finite policy set $\mathcal{D}$ and explicitly normative welfare $W(a, \omega; \nu, h)$, where h is the temporal standpoint, common optimal actions are

$$A^\cap(\Omega) = \bigcap_{\omega \in \Omega} \arg \max_{a \in \mathcal{D}} W(a, \omega; \nu, h).$$

Holding $\mathcal{D}, W, \nu, h$ fixed, define strict uniform robust dominance by $\inf_{\omega \in \Omega} \{W(a, \omega; \nu, h) - W(a', \omega; \nu, h)\} > 0$. If the common-optimum intersection is empty, choice requires an ambiguity rule. Under finite regret suprema, minimax regret is

$$\begin{aligned} \text{Reg}(a; \Omega) &= \sup_{\omega \in \Omega} \left[\max_{a' \in \mathcal{D}} W(a', \omega; \nu, h) - W(a, \omega; \nu, h) \right], \\ A^{\text{MMR}}(\Omega) &= \arg \min_{a \in \mathcal{D}} \text{Reg}(a; \Omega). \end{aligned}$$

Proposition 3 (decision effects of set contraction). If $\emptyset \neq \Omega' \subseteq \Omega$, then $A^\cap(\Omega) \subseteq A^\cap(\Omega')$; strict uniform dominance survives, the undominated set cannot expand, and actionwise and achieved minimax regret weakly fall. Yet minimax-regret action sets need not nest, and every decision object may remain unchanged. A contraction of the supported set is informative, but it has decision value only if it changes a relevant decision object. Bayesian reweighting can matter without shrinking support, while within-set inference may remain prior-sensitive [38]; sampling-aware rules need more structure [39], [40].

For consequential binary transfer $G \in \{0, 1\}$, let m denote mortality assignment, κ beneficiary output per transfer, and v genuine privacy or truthful disclosure to a fixed audience. Normalize disturbance scale to standard logistic and use common intercept -2 :

$$p_{m\kappa v} = P(G = 1 \mid m, \kappa, v) = \Lambda\{-2 + a_m \kappa + b_m + s_m v\}.$$

The scale is normalized; -2 is illustrative, not an identifying normalization. Form and scale remain fixed across arms while coefficients vary: a_m is output sensitivity, b_m impact- and publicity-invariant own-act value, and s_m external-observability sensitivity. None defines an ultimate motive. Own-act value may reflect agency, duty, identity, self-image, relief, demand, or mixtures; publicity may reflect status, shame, accountability, or norm enforcement; output may instrument an actor-side end.

The cue contrast on the normalized index is

$$D_{\kappa v} = \text{logit}(p_{1\kappa v}) - \text{logit}(p_{0\kappa v}) = \Delta a \kappa + \Delta b + \Delta s v.$$

Two impact levels and both publicity states algebraically recover the shifts *within this additive-logit model*. Raw probability interactions do not: nonlinearity can vary a cue effect with publicity even if $\Delta s = 0$. Interpretation requires the constrained law, normalization, additivity, beliefs, and exclusions to survive checks.

Table 1 defines three states sharing one baseline choice law, with $a_0 = 1$ and $b_0 = s_0 = 0$.

Table 1. Three component-shift states with one baseline choice law.

State	Mortality-arm parameters	Model-relative shift
Output	$a_1 = 3, b_1 = s_1 = 0$	beneficiary-output sensitivity
Own-act	$a_1 = 1, b_1 = 2, s_1 = 0$	value of personally transferring
Visibility	$a_1 = 1, b_1 = 0, s_1 = 2$	truthful-observability sensitivity

If only $(\kappa, v) = (1, 1)$ is observed, every state yields

$$p_{0,1,1} = \Lambda(-1) = 0.269, \quad p_{1,1,1} = \Lambda(1) = 0.731.$$

Thus the observed two-arm law is identical. In the mortality arm, compare private high-impact $j_H: (\kappa, v, k^{\text{impl}}) = (2, 0, 0.4)$ with truthful recognition $j_P: (1, 1, 0)$, using net expected beneficiary output:

$$\psi(j, \omega) = \kappa_j p_{1\kappa_j v_j}^\omega - k_j^{\text{impl}},$$

the rankings are shown in Table 2.

Table 2. Behaviorally equivalent baseline effects and policy reversal.

State	$p(j_H)$	$\psi(j_H)$	$p(j_P)$	$\psi(j_P)$	Ranking
Output	0.982	1.564	0.731	0.731	$j_H > j_P$
Own-act	0.881	1.362	0.731	0.731	$j_H > j_P$
Visibility	0.500	0.600	0.731	0.731	$j_P > j_H$

This is an existence example rather than an empirical calibration; it shows that a precise baseline effect can leave a prespecified policy ranking unresolved. Impact or publicity evidence settles robust ranking only when valid exclusions retain states on one side; prior-dependent reweighting may change choice without contracting support. The result is component-relative: impact-sensitive duty or relief may mimic output value, and accountability may mimic image. Adding donor-welfare or legitimacy considerations can enlarge the set of admissible policy rankings.

Tables 1 and 2 distinguish components, not O/S. Separately randomize an actor-return substitute after the cue and before giving. Under the complete-satiation S specification, a non-helping intervention supplies the anticipated relief or meaning restoration of effective giving without providing recipient output; under the restricted O specification, it leaves the beneficiary-relevant content unchanged.

DGPs, O retains index $-c + z\delta\kappa + \varepsilon$, whereas complete-satiation S has

$$-c + z\delta\kappa(1 - Q) + \varepsilon.$$

On the latent binary index, the restricted O specification predicts no cue-by-substitute attenuation, whereas the complete-satiation S specification removes the cue increment.

With valid exclusions and decision-calibrated precision, excluding a prespecified rival region changes a transport prediction: under the complete-satiation S specification, but not under the restricted O specification, a valid actor-return substitute predicts attenuation of mortality-linked giving. It may change policy comparison, although output, experience, menu value, and legitimacy need normative weights. This applies established substitute-return logic to the mortality-reminder setting rather than introducing a new generic contrast.

The audit covers 93 study-level source/version records: all direct assignments, previously display-eligible records, direct economic candidates and unresolved-access risks, key source upgrades, and the 2026 endpoint update. Eighteen records/13 distinct works after consolidating multiple versions pass; 75 fail an ordered gate. Disjoint counts are Tier A: 10 studies/eight works (A1: eight/six; A2: two/two), and Tier B: eight/five. Passing the gates determines inclusion only; it does not imply equal evidentiary weight. The 221-record/121-work map and the complete first-failure ledger are documented in the supplementary audit files. Table 3 reports admitted studies without pooling estimands.

Table 3. Source-qualified evidence, separated by outcome role.

Work; qualifying study (analyzed N)	Assigned cue / identifying feature	Nonstrategic outcome; implementation and recipient	Direct study-level result	Exact source / verification status	Maximum warranted inference
A1 Jonas et al. [4], S2 (22)	own death vs dental pain × domestic/foreign charity	participant-funded charity gift; fully consequential, real recipient	cue×charity on square-root-transformed donations: $F(1,20) = 7.06, p < .02$; raw USD means: domestic 1.44 vs 0.30, foreign 0.88 vs 0.49	Gate 4: I1/I2; paginated version of record (VOR); no preregistration	causal recipient-contingent gift; concern, parochial norm, image, and warm glow unresolved
A1 Ferraro, Shiv, and Bettman [49], S3 (115)	own death vs dental pain; measured virtue	United Way pledge from possible USD 200 prize; lottery-consequential	cue×virtue: $F(1,111) = 6.63, p < .01$; high-virtue raw pledge means USD 65 vs USD 34.50 ($p < .03$)	Gate 4: I1; exact VOR; measured moderator	conditional pledge, not a causal virtue route or altruistic end
A2 Blackie and Cozzolino [5]	own death vs dentist × randomized news claiming high/low blood need	took donor-registration pamphlet; observed low-cost act, no donation	admitted high-need mortality-vs-control coefficient $b = 2.78, p < .05$	Gate 4: I2; checked primary text; death-reflection arm excluded from admitted contrast	response to a represented high-need claim in uptake, not verified need, completed help, or motive
A1 Wade-Benzoni et al. [15], S1–2 (54; 71)	target death vs math news × present/future recipient	charity pledge / energy allocation; S1 lottery-consequential, S2 hypothetical	S1 interaction ($F(1,50)=10.49, p=.002$); S2 ($F(2,65)=3.17, p=.049$)	Gate 4: I1/I2; checked primary full text	future-other contingency; legacy and symbolic self-extension remain live
A2 Xiao, He, and Zhu [6] (240)	own death vs pain × organ/general charity	four-item willingness; hypothetical report	cue ($F(1,236)=14.73, p<.001$); scenario and interaction null	Gate 4: I2; accepted/published r-paginated text	constrains simple rekindled-death avoidance; no revealed gift
A1 Bruine de Bruin and Ulqinaku [7] (5,376)	death-fear items before vs after allocation	USD0–5 self/real-charity allocation; 100 randomly selected decisions implemented	USD 3.65 vs 3.40; $d = .12$, 95% CI [.07, .18]	Gate 4: I1; accepted text; preregistered; version reconciled	small order effect; attention, demand, affect, and motive open
A1 Sætrevik and Sjøstad [38], E2 (781)	own death vs dental pain; time-stamped preregistered close replication of E1's named allocation hypothesis with the same measure in an independent, much larger sample	self/family/friends/charity windfall allocation; hypothetical	study prosocial ratio: mortality 0.177 vs control 0.234 ($p = .036, d = .15$), opposite prediction	Gate 4: I3; exact VOR; preregistered	constrains hypothetical-windfall generalization only

A1 Li and Guan [37], S1–2 (72; 94)	death vs matched negative-affect statements	S1 advantageous-inequity score; S2 equal/self-advantageous choices; one actor payoff lottery-implemented, receiver fictitious	S1: 5.31 vs 3.45 ($d = .549$); S2 at mean-centered self-control: mortality $b = 1.662, p = .020$, interaction $b = -.155, p = .019$	Gate 4: I1; exact open VOR; version collapsed	represented self-benefit and moderated choice, not actual recipient loss, an average effect at every self-control value, a causal self-control route, or an egoistic end
B Salisbury and Nenkov [13], S3–4 (358; 73)	mortality cue and death/live annuity wording	own annuity choice; hypothetical	S3: wording 25.82% vs 35.80% ($p < .05$); essay cue 26.22% vs 34.54% ($p < .09$); interaction $p = .50$. S4: death-worded 26.32% vs live-until 50% ($p < .04$)	Gate 4: I2; checked primary full text	retirement-choice generality only; not in the focal denominator
B Pepper et al. [16], S1–3 (72; 159; 162)	violent-death news; direct preregistrations	own risk and delay choices; hypothetical	all six cue×childhood-SES tests null ($p = .18-.93$)	Gate 4: I3; exact open text; three preregistered replications	failure to reproduce the prespecified cue×childhood-SES risk/time interactions under the adapted violent-crime cue; no other-regard inference
B Williams and James [14] (1,199)	own death vs dental essay × live/death wording	own annuity/bequest choice; hypothetical	mortality ordered-probit coefficient $b = .066, p = .020$; wording main effect null	Gate 4: I2; exact institutional/public text	retirement-choice generality; recipient benefit remains bundled
B Tunney and Raybould [17] (189)	exact Kelley task; preregistered replication	own delay choice; hypothetical	Kelley-style ($p = .088$), opposite trend; MCQ absolute t below 1	Gate 4: I3; exact primary text; preregistered	private-time replication only; no motive conclusion
B Arendsen, van Koningsbruggen, and Fransen [18], S2 (563)	own death vs dentist; preregistered, high-powered test against prespecified $g = .28$ target	status-purchase intention; hypothetical report	3.98 vs 3.93 ($d = .03, p = .746$); source-reported CI internally inconsistent and not used	Gate 4: I3; exact VOR; preregistered	failed to reproduce the prespecified target; not an equivalence test or motive inference

Notes. Tier A (A1/A2) alone is focal; its subtiers are not pooled as implementation-equivalent. Tier B limits success bias and generality but supports no claim about altruism, egoism, giving, helping, or self/other allocation. “Consequential” means a payoff or transfer could occur; actor-only implementation with a fictitious receiver is not prosocial consequence. Each work appears once. Checked-source statistics are transcribed, not re-estimated, harmonized, multiplicity-adjusted, or treated as commensurate weights unless stated. Machine-readable files retain treatment, comparator, delay, recipient reality, implementation, source version, and first failure.

4. Discussion and Conclusion

The present study examined what can legitimately be inferred about motives from mortality-reminder experiments and when additional information about motives becomes economically consequential. The central finding was that random assignment to a mortality reminder can identify a design-specific causal effect on an observed economic choice, but it does not, by itself, identify why that choice changed. In particular, under the paired-construction closure condition developed in this study, specifications in which another person's good has operative noninstrumental standing and specifications in which the actor's own good occupies that role can generate the same complete arm-specific observable distribution. Thus, even when mortality reminders reliably alter giving, allocation, saving, consumption, or another economic choice, the observed treatment effect is insufficient to determine whether the underlying practical content is other-regarding, self-regarding, normative, identity-based, affective, or mixed. This conclusion is consistent with longstanding distinctions between prosocial

behavior and altruistic motivation [19], philosophical accounts emphasizing varieties rather than a single form of altruism [20], and arguments that altruism is a concept whose motivational content cannot be reduced to the material pattern of an action [21]. It is also compatible with Carruthers's account of human motivation, according to which affective benefits and concern for others may coexist without the former automatically displacing the latter as an ultimate motive [22].

This result clarifies why the substantial mortality-salience literature does not support a simple directional claim about human motivation. Earlier meta-analytic evidence suggested that mortality salience can produce reliable effects across psychological and behavioral outcomes [1], but more recent systematic assessment has raised important concerns about publication bias, replicability, and heterogeneity [2]. In the consumer domain, the diversity of mortality-salience effects across products, contexts, mediators, and outcome measures similarly argues against treating mortality awareness as a homogeneous intervention with a single behavioral meaning [3]. The present audit reinforces this position. Although some mortality reminders changed focal economic choices, the high-information studies did not provide a coherent domain-general direction, and the audited evidence did not support a general shift toward either altruistic or self-regarding motivation. This interpretation aligns with preregistered failures to reproduce mortality-salience effects on traditional and novel outcomes [38], failures to reproduce proposed mortality-by-socioeconomic-status interactions in risk and delay preferences [16], and replication evidence indicating that death priming does not consistently alter delay discounting [17]. Recent evidence concerning status consumption likewise indicates that theoretically predicted mortality effects require more qualified interpretation than a general terror-management account would imply [18].

At the same time, the absence of a domain-general motivational conclusion should not be interpreted as evidence that mortality reminders never influence economically relevant behavior. Several studies demonstrate treatment effects under particular institutional and psychological conditions. Mortality salience has been associated with context-dependent changes in charitable behavior [4], prosocial intentions [5, 6], charitable allocation in a national sample [7], and experimentally measured prosocial choices [8]. Mortality-related thought has also been linked to legacy-oriented intergenerational decisions [15], saving behavior [12], retirement-decumulation choices [13], and preferences involving commercial annuities and bequests [14]. More recent work suggests that death reflection and death anxiety may lead to different prosocial pathways [9], while pandemic-related mortality salience may operate through competing processes of social connection and isolation [10]. Endpoint-focused mortality contemplation embedded in mindfulness practice further illustrates that the psychological consequences of thinking about death depend strongly on the form and context of that cognition [11]. Collectively, these studies support the present conclusion that treatment-, institution-, and outcome-specific effects may be meaningful even when they do not justify a common latent-motive interpretation.

The second major finding was that stronger measurement assumptions improve identification but do not automatically recover individual motivational transitions. If an aligned binary component is validly measured in both experimental arms, its treatment and control prevalences identify an arm-marginal effect, and the present analysis showed that these marginals provide sharp Fréchet bounds on the proportion of individuals who could have changed component status. However, they do not identify which particular individuals changed, the direction of all individual changes simultaneously, or the complete joint transition law across multiple nonexclusive components. This result is directly consistent with the logic of partial identification, in which empirical information may restrict a parameter to a set without determining a unique value [40]. It also parallels the broader distinction between likelihood information and assumptions or priors in partially identified models [41]. The practical

implication is that reporting an average increase in an altruism-related measure cannot be reinterpreted as showing that a corresponding proportion of participants became altruistically motivated.

The importance of measurement assumptions is reinforced by recent work on latent causal outcomes. Causal inference involving latent constructs requires an explicit model linking observed indicators to the latent target rather than treating a questionnaire, affective response, or behavioral proxy as the target itself [42]. Nonparametric approaches similarly show that causal identification of latent outcomes depends on substantive measurement and bridge conditions [43]. This finding has particular significance for mortality-reminder research because commonly measured constructs such as emotional satisfaction, death-thought accessibility, identity, intention, or self-reported reasons may be consequences, correlates, mediators, or imperfect indicators of motivation. A behaviorally validated measure of warm glow, for example, can substantially improve construct measurement [29], and experimentally grounded measures can estimate the strength of altruistic motives within a specified model [30]; nevertheless, neither approach makes the mapping from measured construct to ultimate practical content assumption-free.

The distinction is especially evident in warm-glow interpretations. The present results showed that actor-side value can reproduce patterns that might otherwise be attributed to beneficiary-directed concern, particularly when actor return rises with effective helping. This finding accords with analyses questioning whether warm glow is itself an explanation of why people give or instead an emotional consequence or component of giving [27]. Experimental examination of the experienced “warmth” of warm glow likewise demonstrates that affect associated with giving has multiple determinants and does not straightforwardly reveal motivational priority [28]. Zaleskiewicz and colleagues similarly found that mortality salience can increase satisfaction obtained from prosocial behavior [36]; however, the causal-order framework developed here shows why such satisfaction cannot by itself establish that anticipated self-benefit caused the prosocial response. Satisfaction may be sought in advance, may arise because another person's welfare is valued, may be constitutive of successful helping, or may simply follow the choice.

A related finding was that mechanism evidence becomes informative only when its identifying restrictions are explicit. Contemporary causal-methodological research supports this interpretation. Tests of mechanisms can reject sharp hypotheses concerning full mediation or bound portions of treatment effects operating outside a proposed pathway, but rejection of one pathway does not uniquely establish another [44]. Likewise, tests of full mediation and causal-mechanism identifiability demonstrate that mediation evidence remains conditional on assumptions about the causal structure and observability of the mechanism [45]. In the current framework, this is captured by the four bridges of content, choice, measurement, and causal order. A mortality cue may alter the construal of an action rather than simply change the weight placed on a fixed motive, which is consistent with reason-based accounts in which decision-relevant properties and reasons vary with context [35]. The finding that self-control can moderate mortality-related fairness decisions further supports the idea that the same cue can operate through heterogeneous decision processes [37].

The third major result concerned the economic decision value of motive evidence. The analysis demonstrated that narrowing the set of motive-consistent states is scientifically informative but does not necessarily alter a decision. Motive information has economic value when it changes a prespecified prediction, transport conclusion, robust-dominance relation, minimax-regret choice, policy ranking, or welfare assessment. This follows naturally from contemporary decision theory under partial identification, where policy selection must explicitly accommodate unresolved uncertainty about payoff-relevant states [47]. Recent work on decision rules with partially identified payoffs similarly demonstrates that optimal action need not require complete point

identification but does require a transparent relation between the supported state set, losses, and available actions [48]. The finding also complements behavioral-welfare arguments that observed choices should not automatically be treated as welfare-revealing when contextual or psychological forces may distort the relation between choice and welfare [46]. Attempts to “purify” preferences therefore require substantive assumptions about which observed influences are normatively relevant rather than simply statistical adjustment [34].

This decision-theoretic point was illustrated by the finding that states producing the same baseline mortality-reminder choice effect could reverse the ranking of beneficiary-impact and public-recognition policies. When the treatment effect reflected beneficiary-output sensitivity, increasing actual beneficiary impact was favored; when it reflected observability, public recognition could instead become the effective behavioral margin. This is closely related to economic work showing that carefully chosen experimental variation can distinguish components of giving behavior. Variation in outside provision can separate public-output sensitivity from own-gift value within maintained models [24], and comparative work across economics and psychology emphasizes that “altruism,” “egoism,” and “warm glow” often refer to different experimental targets [25]. More recent synthesis similarly advocates unifying motive research through clearer distinctions between behavioral components and psychological interpretations [26]. The present findings extend this logic to mortality reminders by showing that knowing that donations rose is insufficient to determine which intervention should follow from that result.

The policy illustration also aligns with research manipulating other diagnostic margins. DellaVigna and colleagues showed that advance notice and opportunities to avoid solicitation can distinguish giving value from social-pressure effects [31], demonstrating why an observed donation cannot automatically be interpreted as welfare-improving altruism. Research on moral-preference elicitation under image concerns similarly shows that observability can systematically alter expressed moral choices and must therefore be treated as a distinct behavioral component [32]. Intertemporal altruism research demonstrates that separating the timing of choice from the timing of beneficiary consequences can identify distinct components of prosocial utility [33]. These studies support the current argument that informative experimental contrasts should cross-cut rival explanations rather than simply add more measures of the same behavioral response. They also explain why the proposed mortality-cue $\times$ beneficiary-impact $\times$ publicity design could distinguish model-relative components while still falling short of identifying unrestricted ultimate motives.

Finally, the evidence map showed why synthesis should occur at the level of warranted inference rather than through simple pooling. Only a subset of audited studies met the complete source-sensitive criteria, and even those studies differed in cue, comparator, consequence structure, recipient reality, outcome, and estimand. This heterogeneity makes vote counting or a domain-general pooled interpretation inappropriate. Evidential triangulation is strongest when methods with genuinely different weaknesses converge on the same well-defined target [39]; merely combining heterogeneous outcomes does not produce such triangulation. Likewise, mixed-motive theories of giving already predict that similar actions can emerge from multiple combinations of concern, norms, reciprocity, and self-regard [23]. The findings therefore suggest that mortality-reminder research should shift from asking whether death reminders make people generically “more altruistic” or “more selfish” toward specifying the precise economic choice component, practical-content target, identifying bridges, and decision problem for which the evidence is intended.

Several limitations should be considered when interpreting these findings. The evidence map was deliberately search-bounded rather than exhaustive, included machine-assisted preliminary screening, was affected by retrieval caps and unresolved source access, and was not independently duplicated across all subscription-based databases.

The source-qualified studies were highly heterogeneous in experimental cue, comparator, population, timing, consequentiality, observability, outcome construction, and estimand, which prevented meaningful statistical pooling. The formal conclusions also depend on the specified admissible classes, alignment rules, and bridge assumptions; changing these assumptions can widen or narrow the identified set. Moreover, arm-level data do not provide the within-person coupling needed to identify complete individual motive transitions, and self-report or process measures remain vulnerable to reactivity, shared measurement error, and uncertainty about causal order. Finally, the proposed experimental designs were analytical implications of the identification framework and were not implemented within the present study.

Future studies should design mortality-reminder experiments around explicitly competing predictions rather than adding additional correlational motive measures after observing a treatment effect. A particularly useful approach would be to preregister factorial experiments that independently manipulate mortality cues, verified beneficiary impact, and truthful observability while holding beneficiary identity, budget, information, and timing constant. Equivalence regions should be defined in advance and calibrated to the policy decision that the experiment is intended to inform. Subsequent experiments could introduce independently validated actor-return substitutes to distinguish beneficiary-sensitive responses from anticipated relief, meaning restoration, or other actor-side returns. Repeated-measure or crossover designs may also help narrow transition uncertainty when stability, carryover, and content-alignment assumptions can be defended. Future evidence syntheses should use independent screening, complete source reconciliation, study-level estimand coding, and explicit separation of consequential behavior, hypothetical intention, strategic interaction, and private-choice outcomes.

Practitioners should avoid interpreting an increase in donations, helping, saving, or other mortality-linked behavior as direct evidence that mortality reminders have increased altruism, selfishness, generosity, or any other broad motive. Fundraising and consumer interventions should instead be selected according to the behavioral component that has actually been identified. If responses reliably track verified beneficiary impact, improving effectiveness and communicating that effectiveness may be preferable to increasing visibility; if they track observability, recognition may be behaviorally relevant but should be assessed separately for legitimacy and participant welfare; and if advance notice or private refusal substantially changes behavior, pressure and autonomy should be included in the decision criterion. Mortality-related interventions should also include safeguards for distress, preserve meaningful opportunities to decline exposure or participation, and avoid using behavioral responsiveness as evidence that the intervention benefits the person being influenced.

Authors' Contributions

Authors equally contributed to this article.

Ethical Considerations

All procedures performed in this study were under the ethical standards.

Acknowledgments

Authors thank all participants who participate in this study.

Conflict of Interest

The authors report no conflict of interest.

Funding/Financial Support

According to the authors, this article has no financial support.

References

- [1] B. L. Burke, A. Martens, and E. H. Faucher, "Two Decades of Terror Management Theory: A Meta-Analysis of Mortality Salience Research," *Personality and Social Psychology Review*, vol. 14, no. 2, pp. 155-195, 2010, doi: 10.1177/1088868309352321.
- [2] L. Chen, R. Benjamin, Y. Guo, A. Lai, and S. J. Heine, "Managing the Terror of Publication Bias: A Systematic Review of the Mortality Salience Hypothesis," *Journal of Personality and Social Psychology*, vol. 129, no. 1, pp. 20-41, 2025, doi: 10.1037/pspa0000438.
- [3] J. L. Arendsen, B. C. Brugman, G. M. van Koningsbruggen, and M. L. Fransen, "Terror Management Theory in the Consumer Domain: A Systematic Review and Meta-Analysis on Mortality Salience Driven Consumer Responses," *International Journal of Consumer Studies*, vol. 49, no. 3, p. e70056, 2025, doi: 10.1111/ijcs.70056.
- [4] E. Jonas, D. Sullivan, and J. Greenberg, "Generosity, Greed, Norms, and Death—Differential Effects of Mortality Salience on Charitable Behavior," *Journal of Economic Psychology*, vol. 35, pp. 47-57, 2013, doi: 10.1016/j.joep.2012.12.005.
- [5] L. E. R. Blackie and P. J. Cozzolino, "Of Blood and Death: A Test of Dual-Existential Systems in the Context of Prosocial Intentions," *Psychological Science*, vol. 22, no. 8, pp. 998-1000, 2011, doi: 10.1177/0956797611415542.
- [6] Q. Xiao, W. He, and Y. Zhu, "Re-Examining the Relationship between Mortality Salience and Prosocial Behavior in Chinese Context," *Death Studies*, vol. 41, no. 4, pp. 251-255, 2017, doi: 10.1080/07481187.2016.1268220.
- [7] W. Bruine de Bruin and A. Ulqinaku, "Effect of Mortality Salience on Charitable Donations: Evidence from a National Sample," *Psychology and Aging*, vol. 36, no. 4, pp. 415-420, 2021, doi: 10.1037/pag0000478.
- [8] T. Bao, X. Li, and C. Xia, "Life Is Too Short to Be Small: An Experiment on Mortality Salience and Prosocial Behavior," *China Economic Review*, vol. 86, p. 102191, 2024, doi: 10.1016/j.chieco.2024.102191.
- [9] T. Y. Li, Y. H. Mao, X. Y. Liang, C. M. Jiang, and H. Y. Sun, "Death Reflection vs. Death Anxiety: Divergent Paths to Pro-Sociality," *Death Studies*, pp. 1-12, 2025, doi: 10.1080/07481187.2025.2607432.
- [10] L. Meng *et al.*, "Connecting or Isolating: Investigating the Influence of Pandemic Mortality Salience on Prosocial Intention," *Acta Psychologica Sinica*, vol. 57, no. 4, pp. 652-670, 2025, doi: 10.3724/SP.J.1041.2025.0652.
- [11] F. Xiao, Q. Peng, and X. T. Wang, "Mindfulness with Mental Time Travel: Emotional and Behavioral Effects of Endpoint-Focused Meditation," *Mindfulness*, 2026, doi: 10.1007/s12671-026-02959-8.
- [12] T. Zaleskiewicz, A. Gasiorowska, and P. Kesebir, "Saving Can Save from Death Anxiety: Mortality Salience and Financial Decision-Making," *PLoS One*, vol. 8, no. 11, p. e79407, 2013, doi: 10.1371/journal.pone.0079407.
- [13] L. C. Salisbury and G. Y. Nenkov, "Solving the Annuity Puzzle: The Role of Mortality Salience in Retirement Savings Decumulation Decisions," *Journal of Consumer Psychology*, vol. 26, no. 3, pp. 417-425, 2016, doi: 10.1016/j.jcps.2015.10.001.
- [14] J. A. Williams and R. N. James, "Bequest Provision Preferences in Commercial Annuities: An Experimental Test of the Role of Mortality Salience," *Journal of Financial Counseling and Planning*, vol. 30, no. 1, pp. 121-131, 2019, doi: 10.1891/1052-3073.30.1.121.
- [15] K. A. Wade-Benzoni, L. Plunkett Tost, M. Hernandez, and R. P. Larrick, "It's Only a Matter of Time: Death, Legacies, and Intergenerational Decisions," *Psychological Science*, vol. 23, no. 7, pp. 704-709, 2012, doi: 10.1177/0956797612443967.
- [16] G. V. Pepper, D. H. Corby, R. Bamber, H. Smith, N. Wong, and D. Nettle, "The Influence of Mortality and Socioeconomic Status on Risk and Delayed Rewards: A Replication with British Participants," *PeerJ*, vol. 5, p. e3580, 2017, doi: 10.7717/peerj.3580.
- [17] R. J. Tunney and J. N. Raybould, "Thinking about Neither Death nor Poverty Affects Delay Discounting, but Episodic Foresight Does: Three Replications of the Effects of Priming on Time Preferences," *Quarterly Journal of Experimental Psychology*, vol. 76, no. 4, pp. 838-849, 2023, doi: 10.1177/17470218221097047.
- [18] J. Arendsen, G. M. van Koningsbruggen, and M. L. Fransen, "Mortality Salience Effects and Status Consumption—A Reasoned Action Approach," *PLoS One*, vol. 21, no. 5, p. e0350005, 2026, doi: 10.1371/journal.pone.0350005.
- [19] S. Pfattheicher, Y. A. Nielsen, and I. Thielmann, "Prosocial Behavior and Altruism: A Review of Concepts and Definitions," *Current Opinion in Psychology*, vol. 44, pp. 124-129, 2022, doi: 10.1016/j.copsyc.2021.08.021.
- [20] P. Kitcher, "Varieties of Altruism," *Economics and Philosophy*, vol. 26, no. 2, pp. 121-148, 2010, doi: 10.1017/S0266267110000167.
- [21] M. Schefczyk and M. Peacock, "Altruism as a Thick Concept," *Economics and Philosophy*, vol. 26, no. 2, pp. 165-187, 2010, doi: 10.1017/S0266267110000180.

- [22] P. Carruthers, *Human Motives: Hedonism, Altruism, and the Science of Affect*. Oxford, UK: Oxford University Press, 2024.
- [23] J. Konow, "Mixed Feelings: Theories of and Evidence on Giving," *Journal of Public Economics*, vol. 94, no. 3-4, pp. 279-297, 2010, doi: 10.1016/j.jpubeco.2009.11.008.
- [24] M. Ottoni-Wilhelm, "Altruism and Egoism/Warm Glow in Economics and Psychology: Building a Bridge between Different Experimental Approaches," in *Social Economics: Current and Emerging Avenues*, J. Costa-Font and M. Macis Eds. Cambridge, MA, USA: MIT Press, 2017, pp. 79-106.
- [25] M. Ottoni-Wilhelm, L. Vesterlund, and H. Xie, "Why Do People Give? Testing Pure and Impure Altruism," *American Economic Review*, vol. 107, no. 11, pp. 3617-3633, 2017, doi: 10.1257/aer.20141222.
- [26] M. Ottoni-Wilhelm and L. Vesterlund, "Motives to Give in Economics and Psychology: A Step toward Unification," *Journal of Economic Behavior & Organization*, vol. 216, pp. 354-365, 2023, doi: 10.1016/j.jebo.2023.09.018.
- [27] R. T. Bianchi, "Warm-Glow Effects and Warm-Glow Explanations: Does the Warm Glow from Giving Provide an Answer to the Question of Why People Give?," University of Geneva, 2022. [Online]. Available: <https://archive-ouverte.unige.ch/unige:158924>
- [28] R. Bianchi, F. Cova, and E. Tieffenbach, "Is the Warm Glow Actually Warm? An Experimental Investigation into the Nature and Determinants of Warm Glow Feelings," *International Journal of Wellbeing*, vol. 13, no. 3, pp. 1-23, 2023, doi: 10.5502/ijw.v13i3.2565.
- [29] J. Carpenter, A. Lyford, and M. Zhang, "A Behaviorally Validated Warm Glow Questionnaire," *Journal of the Economic Science Association*, vol. 10, no. 2, pp. 310-329, 2024, doi: 10.1007/s40881-024-00161-x.
- [30] N. W. Chan, S. Knowles, R. Peeters, and L. Wolk, "Measuring Strength of Altruistic Motives," *Journal of the Economic Science Association*, vol. 10, no. 2, pp. 595-602, 2024, doi: 10.1007/s40881-024-00170-w.
- [31] S. DellaVigna, J. A. List, and U. Malmendier, "Testing for Altruism and Social Pressure in Charitable Giving," *Quarterly Journal of Economics*, vol. 127, no. 1, pp. 1-56, 2012, doi: 10.1093/qje/qjr050.
- [32] R. Bénabou, A. Falk, L. Henkel, and J. Tirole, "Eliciting Moral Preferences under Image Concerns: Theory and Experiment," 2026. [Online]. Available: https://luca-henkel.github.io/papers/Moral_Preferences.pdf
- [33] F. Chopra, A. Falk, and T. Graeber, "Intertemporal Altruism," *American Economic Journal: Microeconomics*, vol. 16, no. 1, pp. 329-357, 2024, doi: 10.1257/mic.20210319.
- [34] L. Beck, "The Econ Within or the Econ Above? On the Plausibility of Preference Purification," *Economics and Philosophy*, vol. 39, no. 3, pp. 423-445, 2023, doi: 10.1017/S0266267122000141.
- [35] F. Dietrich and C. List, "Reason-Based Choice and Context-Dependence: An Explanatory Framework," *Economics and Philosophy*, vol. 32, no. 2, pp. 175-229, 2016, doi: 10.1017/S0266267115000474.
- [36] T. Zaleskiewicz, A. Gasiiorowska, and P. Kesebir, "The Scrooge Effect Revisited: Mortality Salience Increases the Satisfaction Derived from Prosocial Behavior," *Journal of Experimental Social Psychology*, vol. 59, pp. 67-76, 2015, doi: 10.1016/j.jesp.2015.03.005.
- [37] W. Li and L. Guan, "Self-Control Buffers the Mortality Salience Effect on Fairness-Related Decision-Making," *Behavioral Sciences*, vol. 14, no. 12, p. 1121, 2024, doi: 10.3390/bs14121121.
- [38] B. Sætrevik and H. Sjøstad, "Mortality Salience Effects Fail to Replicate in Traditional and Novel Measures," *Meta-Psychology*, vol. 6, p. MP.2020.2628, 2022, doi: 10.15626/MP.2020.2628.
- [39] J. Kuorikoski and C. Marchionni, "Evidential Diversity and the Triangulation of Phenomena," *Philosophy of Science*, vol. 83, no. 2, pp. 227-247, 2016, doi: 10.1086/684960.
- [40] E. Tamer, "Partial Identification in Econometrics," *Annual Review of Economics*, vol. 2, pp. 167-195, 2010, doi: 10.1146/annurev.economics.050708.143401.
- [41] H. R. Moon and F. Schorfheide, "Bayesian and Frequentist Inference in Partially Identified Models," *Econometrica*, vol. 80, no. 2, pp. 755-782, 2012, doi: 10.3982/ECTA8360.
- [42] L. F. Stoetzer, X. Zhou, and M. Steenbergen, "Causal Inference with Latent Outcomes," *American Journal of Political Science*, vol. 69, no. 2, pp. 624-640, 2025, doi: 10.1111/ajps.12871.
- [43] J. Fu and D. P. Green, "Nonparametric Identification and Estimation of Causal Effects on Latent Outcomes," arXiv, 2026. [Online]. Available: <https://doi.org/10.48550/arXiv.2604.08681>
- [44] S. Kwon and J. Roth, "Testing Mechanisms," *Review of Economic Studies*, p. rdag028, 2026, doi: 10.1093/restud/rdag028.
- [45] M. Huber, K. Kloiber, and L. Laffers, "Testing Full Mediation of Treatment Effects and the Identifiability of Causal Mechanisms," arXiv, 2026. [Online]. Available: <https://doi.org/10.48550/arXiv.2603.04109>
- [46] B. D. Bernheim and A. Rangel, "Beyond Revealed Preference: Choice-Theoretic Foundations for Behavioral Welfare Economics," *Quarterly Journal of Economics*, vol. 124, no. 1, pp. 51-104, 2009, doi: 10.1162/qjec.2009.124.1.51.
- [47] J. L. Montiel Olea, C. Qiu, and J. Stoye, "Decision Theory for Treatment Choice Problems with Partial Identification," *Review of Economic Studies*, p. rdag015, 2026, doi: 10.1093/restud/rdag015.
- [48] T. Christensen, H. R. Moon, and F. Schorfheide, "Optimal Decision Rules When Payoffs Are Partially Identified," *Review of Economic Studies*, p. rdag017, 2026, doi: 10.1093/restud/rdag017.

- [49] R. Ferraro, B. Shiv, and J. R. Bettman, "Let Us Eat and Drink, for Tomorrow We Shall Die: Effects of Mortality Salience and Self-Esteem on Self-Regulation in Consumer Choice," *Journal of Consumer Research*, vol. 32, no. 1, pp. 65-75, 2005, doi: 10.1086/429601.