

Design Variety Measurement using Sharma-Mittal Entropy

Faez Ahmed*

Dept. of Mechanical Engineering
Northwestern University
Evanston, Illinois 60208
Email: faez@northwestern.edu

Sharath Kumar Ramachandran

School of Engineering Design, Technology
and Professional Programs
The Pennsylvania State University
University Park, PA
Email: sharath@psu.edu

Mark Fuge

Dept. of Mechanical Engineering
University of Maryland
College Park, Maryland 20742
Email: fuge@umd.edu

Sam Hunter

Industrial and Organizational Psychology
The Pennsylvania State University
University Park, PA
Email: sth11@psu.edu

Scarlett Miller

School of Engineering Design,
Technology and Professional Programs
The Pennsylvania State University
University Park, PA
Email: shm13@psu.edu

Design variety metrics measure how much a design space is explored. This paper proposes that a generalized class of entropy metrics based on Sharma-Mittal entropy offers advantages over existing methods to measure design variety. We show that an exemplar metric from Sharma-Mittal entropy, named the Herfindahl–Hirschman Index for Design (HHID) has the following desirable advantages over existing metrics: (a) More Accuracy: It better aligns with human ratings compared to existing and commonly used tree-based metrics for two new datasets; (b) Higher Sensitivity: It has higher sensitivity compared to existing methods when distinguishing between the variety of sets; (c) Allows Efficient Optimization: It is a submodular function, which enables one to optimize design variety using a polynomial-time greedy algorithm; and (d) Generalizes to Multiple Metrics: Many existing metrics can be derived by changing the parameters of this metric, which allows a researcher to fit the metric to better represent variety for new domains. The paper also contributes a procedure for comparing metrics used to measure variety via constructing ground truth datasets from pairwise comparisons. Overall, our results shed light on some qualities that good design variety metrics should possess and the non-trivial challenges associated with collecting the data needed to measure those qualities.

NOMENCLATURE

SVS Design variety metric proposed in Shah *et al.* [1]
NM Design variety metric proposed by Nelson *et al.* [2]
HHI Herfindahl–Hirschman Index [3]
SME Sharma-Mittal Entropy [4]

INTRODUCTION

Creativity is the capacity to generate unique and original work that is useful [5–7]. Creative solutions help individuals in solving day-to-day tasks and societies by yielding meaningful scientific findings [6].

Past research [8] relates creativity with divergent thinking — the capacity to produce a wider variety of ideas with higher fluency. Divergent thinking has been shown to correlate with the success of the final product [9–12]. Prior work supports that chances of solving a problem increase when a more diverse set of ideas is produced in the initial stages of the design process [1, 13, 14]. These findings encourage the need to explore the design space in the early stages of design [15]. But how does one quantify design space exploration?

Engineering researchers have sought to capture how “explored the solution space” is by measuring design variety [1]. There are two approaches typically deployed in

*Address all correspondence to this author.

engineering literature to measure design variety: qualitative (or subjective) and quantitative ratings of variety. As one example of subjectively evaluating design variety, Linsey *et al.* [16] proposed taking a set of ideas and dividing them into pools based on intuitive categories created by the coder. At the completion of this sorting process, an individuals variety score is determined by counting the number of bins into which their ideas were sorted and dividing that number by the total number of bins. The metric relies on a rater's mental model rather than a quantitative procedure. While these subjective ratings provide a relatively efficient method for measuring design variety in terms of the amount of time and effort required to code design variety, this efficiency comes at the potential cost of the validity and reliability of the metric [17]. This paper does not investigate this class of subjective variety metrics.

In contrast to subjective ratings, the other approach to measure design variety is using a quantitative approach. Within quantitative approaches, genealogical tree approaches are widely used to measure variety, as evident by hundreds of studies citing these approaches [1, 2, 18]. In these approaches, subjective human raters are replaced with a deterministic formula that depends on a few measured attributes of a set of designs. One of the first metrics to use this approach was developed by Shah, Smith and, Vargas-Hernandez [1] (SVS metric) who broke each design into four hierarchical levels (physical principle, working principle, embodiment, and detail) to calculate design variety. The SVS metric is repeatable and attempts to reduce subjectivity by using predefined criteria for measuring variety. However, researchers have reported a lack of sensitivity and accuracy of SVS [19–21]. For example, the genealogical tree calculation method (like SVS) is inconsistent with experts ratings of variety [19]. Besides, studies have shown that the sensitivity of the SVS metric diminishes when it is applied to large datasets [20] due to the exclusion of important abstract differences and generally focuses on dissimilarity in the embodiment level [21].

While computing a variety metric score for a set of ideas is straightforward, finding which subset of ideas has the highest variety score often requires computing the score for all possible combinations. For large datasets, this process becomes computationally expensive for any variety metric which does not have a computationally tractable method of optimization (for example, more than one billion SVS evaluations would be needed for finding six out of hundred ideas with the highest SVS score). Finally, different metrics may be more suitable for different domains. However, there is a lack of understanding of connections between these metrics, which are measuring the same underlying phenomenon of variety in a domain. While Fuge *et al.* [22] argued that a few variety metrics belong to a family of mathematical functions, they did not propose a single parametrized function which can unify many variety measurement methods.

This paper re-examines two of these hierarchical metrics and compare them to methods of calculating diversity from other (non-engineering) domains. Specifically, this paper compares the tree-based metrics of SVS [1] and NM [2] with

the entropy-based measure of the Herfindahl–Hirschman Index (HHI). This paper shows how to adapt HHI to Engineering Design problems and proposes a new metric named Herfindahl–Hirschman Index for Design (HHID). By comparing HHID to SVS [1] and NM [2], this paper argues and empirically demonstrates that HHID is a more accurate and sensitive measure for variety that has clear benefits for engineering and design measurement applications. We also demonstrate that this new metric is optimizable and can be generalized using a broader class of metrics.

The key contributions of this paper can be categorized under five themes — Accuracy, Sensitivity, Optimizability, Generalizability, and a Ground-truth Dataset:

1. **Accuracy:** This paper proposes a measurement procedure that can estimate the accuracy of variety metrics via alignment with ground truth datasets comprising of pairwise comparisons. A ground truth dataset refers to a set of design examples, where one is confident of the variety measurement. Results 2 and 7 discuss how they are established using pairwise comparisons. Any metric which gives the same relative scores as a query in this dataset is considered accurate. Using a new family of variety metrics, the findings indicate that entropy-based metrics better align with human judgments of variety compared to two existing tree-based metrics for two datasets used in this study.
2. **Sensitivity:** This paper proposes a method of approximating metric sensitivity by randomly selecting sets and comparing their scores. The analysis shows that the SVS and NM metrics give the same variety score to a large percentage of sets (approximately 30% for our dataset), while HHID index has higher sensitivity in distinguishing between different sets of ideas.
3. **Optimizability:** The metric functions proposed in this paper are monotone non-decreasing and submodular¹, which allows one to propose a scalable greedy optimization algorithm with a constant factor optimality guarantee. To find a set of five designs with the highest variety from a collection of 1000 designs, brute force using traditional metrics makes more than 8 trillion metric evaluations, while greedy optimization gives the near-optimal solution in less than 5000 metric evaluations. This represents an efficiency improvement of around six orders of magnitude for even just a modest sized problem.
4. **Generalizability:** The paper proposes that a general class of entropy-based metrics based on Sharma-Mittal entropy can be used to measure variety. We discuss how the choice of two parameters in the Sharma-Mittal entropy family affects the type of variety one wants to measure. This enables one to customize the behavior of the variety metric to a broader set of behaviors that current variety metrics can model.
5. **Ground-truth Dataset:** The study leads to two datasets of pairwise comparisons which are released for future researchers to use. These datasets with pairwise queries

¹Submodular functions are set functions to model diminishing marginal utility.

can be used as a common scale to measure improvement in variety metrics in future studies.

The proposed family of metric has a few limitations. First, our experiments were limited to two datasets. Second, when different attributes have hierarchical relationships (*e.g.* an electric motor is dependent on electricity as the mode of power), the proposed metrics do not model these relationships while calculating variety.

BACKGROUND AND RELATED WORK

This section first reviews some qualities that good variety metrics should possess. Then, it discusses what factors researchers should consider when constructing a ground truth evaluation method for comparing variety metrics.² Lastly, it reviews existing design variety metric literature.

What Qualities Should a Good Design Metric Possess?

Quality control is essential when creating and evaluating metrics that map abstract concepts like creativity to quantitative metrics. Particularly when metrics can be either subjective or objective, researchers need to demonstrate that they are valid and reliable without circularity [23]. Design metrics can be relative or absolute. Relative design metrics compare ideas against other ideas in the same generated set [24]. In this way, designs generated in the same design session addressing the same problem can be compared and contrasted to tease out designs to develop further. In contrast, absolute metrics are not dependent on what other ideas are in the set. Researchers also need to reduce the subjectivity in measurement techniques, so the results do not depend on individual judges. For example, in the field of psychometrics, researchers try to craft sets of questions that produce internally consistent results — that is, if one asks the same questions one should get repeatable, similar answers even under minor changes to the test environment or experimental setup [25]. However, these questions only ensure repeatability and not validity. Validity refers to the extent to which a measurement reflects the absolute state of an artifact under observation — that is, a ground truth. The term “valid” refers to an external frame of reference or a universally accepted standard against which a measurement is tested [26]. Many creativity metrics leverage a rater’s expertise in a given domain to ensure metric validity [27]. This is necessary to eliminate circularity or measuring un-validated metrics against other un-validated metrics [28]. In this paper, the term “accuracy” is used to measure the validity of a metric against a known standard.

The key assumption in many past works is that raters who have considerable experience in a given domain are best suited to provide ground truth assessment for tasks like evaluating creativity [29]. If experts are the de-facto ground truth, then why do we need a separate, objective metric? This is because of resource and practical constraints: expert time and effort is a scarce commodity. This scarcity forces

researchers to develop objective metrics that can aid quasi-experts or novice raters in accurately evaluating processes and ideas.

But how do we verify whether a proposed metric is valid, or, equivalently, accurately measures creativity? This paper focuses on how to validate any proposed objective metric against expert raters. The paper focuses on how variety metrics must be evaluated to ensure they are measuring what they are built to measure, reliably, and with an acceptable degree of validity.

When a metric is created, it is important to establish some desiderata (qualities we want) that a metric must possess. Prior work on establishing acceptable qualities of a metric includes the work of Simonton and Amabile [30], who were key in standardizing the measurement of creativity in psychological research. Previously, most methods used pencil and paper tests, personality tests, biographical inventories (such as Schaefer and Anastasi’s biographical inventory [31] and Taylor’s Alpha Biographical Inventory [32]) *etc.* These tests were debatable in experiments that sought to reduce within-group variability and generally lacked a clear creativity definition and an effective strategy to avoid biases on behalf of the rater [30]. This seminal work highlighted the need to better understand the multiple desiderata for a creativity metric. Building upon that work, this paper attempts to mathematically describe and lay out experimental procedures by which one might measure such desiderata.

For example, good metrics should have the ability to establish ground truths using expert agreements and must be replicable by other raters who use the metric. In this regard, variety metrics like SVS and NM were developed to reduce subjectivity on the rater’s part and make it easier for researchers to replicate a processes used to analyze designs. For subjective metrics, high inter-rater reliability and internal consistency are reported as some of the desired qualities of the metric [33].

This paper argues that for any new design variety metric, accuracy, sensitivity, optimizability, repeatability, and explainability are also desirable qualities. Here accuracy means that if ground truth estimates of a quantity are available, then a new metric should align with this ground truth. Sensitivity means that a new metric should be able to distinguish changes between different states of a quantity. Repeatability means that when measurements are repeated again and again with the quantity being unchanged, they should not give different measurements. Explainability means that the measuring instrument should give explainable scores, that is, it should be possible to explain why one set of designs received a higher score than another set. Finally, optimizability means that given a ground set of ideas, if the goal is to find sets of ideas that will have maximum or minimum measurement score using a variety metric, then practitioners should be able to do so in polynomial time (where time is a simple polynomial function of the size of the input—for example the number of designs considered). Subjective metrics generally lack repeatability and explainability. In contrast, existing metrics like SVS and NM are repeatable and explainable. However, this work shows that SVS and NM are not

²We use ‘variety metrics’ for metrics used to measure design variety.

accurate, sensitive and optimizable compared to the Sharma-Mittal family of metrics.

Why is Measuring Design Variety Important?

Engineering researchers introduced design variety metrics to measure how well someone explores the solution space during a design task [34]. Generating a large number of ideas with iterative or small changes may not result in effective concept generation or innovative products. Research has shown that “there is no way to generate an optimum solution without exploring the solution space through early tentative ideas” (Pg. 11 [35]), which shows the importance of measuring design variety. Hence, the potential to develop ideas of broad variety is correlated with the ability to successfully re-construct and solve problems. This ability is referred to as cognitive restructuring in psychology [1]. Cognitive restructuring is frequently used along in concert with the number of ideas developed (quantity) to assess design ideation.

Research in Engineering Design has shown a correlation between the amount of design space explored and the quality of the final design [11]. Without exploration, designers may misconstrue the solution space to be narrow [36]. One of the main contributing factors to this trend is functional fixation, or blind adherence to solutions that are familiar and comfortable, which can generally lead to products of lower quality or innovation [37,38]. Suppose you have ten teams, each of which generates five designs. Your task is to select one of these sets and use it as an inspiration for new designs. Ideally, you would want to select a set which causes minimum functional fixation and is also high quality. How should one do this? To choose this set and know how much a design space is explored, one needs to measure both the quality of the designs and their variety. This work focuses on the measurement of design variety.

Measuring design space exploration requires computing mathematical functions on groups of ideas [17]. To address this need to measure the extent to which tools promote variety, Shah *et al.* [1] developed a metric (SVS) to provide a repeatable and reliable method to calculate design variety by rewarding ideas that are differentiated at higher levels of abstraction. In the SVS metric, the authors decompose design variety into four hierarchical levels: the physical principle, followed by the working principle, embodiment, and detail. Shah *et al.* proposed that design variety should be calculated as shown below in equation 1.

$$V = \sum_{j=1}^m (f_j) \sum_{k=1}^4 (S_k \cdot B_k) / N \quad (1)$$

where V is the variety score, m is the number of functions solved by the design, f_j is a weight assigned to the relative importance of function j , S_k is the score for hierarchical level k , B_k is the number of branches at hierarchical level k , and N is the total number of ideas in the set. The key intuition behind this metric is that a set of ideas can be represented by

a tree comprised of hierarchical attributes. Attributes on top of the hierarchy are more important than ones below, and if a set has multiple ideas with unique higher-level attributes, then that set gets a higher variety score.

Researchers have found that the SVS metric double counts ideas at each level in the tree and there is lack of guidance on how the specific numerical choice of the weights at each level of the tree is to be determined [2, 39]. Because of these pitfalls, Nelson *et al.* [2] refined the metric by seeking to account for the double-counting of ideas present in the SVS metric by considering the number of differentiation at each hierarchical level rather than considering all the levels. Besides, Nelson *et al.* modified the SVS metric by altering the weighting scheme from 10, 6, 3 & 1 to 10, 5, 2 & 1 for the physical principle, the working principle, the embodiment, and detail respectively. They argued that the new weighting scheme assures that at least two ideas at a lower hierarchical level must be added to equal the variety gain by adding a single idea at the next higher hierarchical level [2].

However, both SVS and NM do not define what each level of the hierarchy means. There has been insufficient empirical justification or verification of the weights used in such genealogical tree metrics [19], which can lead to large variations in the deployment of the metric in Engineering Design research. Srinivasan and Chakrabarti [40] also propose an idea space variety metric and base it on the abstraction levels of the SAPPPhIRE model, with different weights for action, state change, input, phenomenon, effect, organ and part abstraction levels. Other improvements of SVS metric includes the work of Verhaegen *et al.* [18], who combined Shah’s metric with a Herfindahl-index-based tree entropy penalty, to encourage “uniformness of distribution” — essentially preferring trees that have even branching. In Ref. [41], the authors showed problems with many tree-based metrics, including arbitrarily defined weights. However, most of these methods require constructing a hierarchical tree. Our analysis demonstrates that the additional step of constructing a tree may not be necessary for measuring design variety as seemingly simple entropy metrics (which do not require tree construction) often perform better.

Apart from using hierarchical trees, researchers have employed many other variety metrics of varying complexity. A simple measure of variety is the ratio of the number of categories that a participants’ ideas occupied to the total number of bins, which has been used to understand design fixation [16,42]. Later, it is shown that a generalized two-parameter metric reduces to form very similar to the above metric by selecting both parameters to zero. In Ref. [43], the authors proposed a Comprehensive Metric for Commonality (CMC) to evaluate the design of a product family based on product attributes and the allowed diversity in the family. Henderson *et al.* compared different variety metrics and propose a new metric which calculates variety by looking at how a collection of ideas covers a potential design space based on the diversity of the other metrics used to assess those ideas. While most metrics discussed so far focused on measuring variety in an ideation exercise, Kota *et al.* [44] proposed Product Line Commonality Index (PCI), to capture the level

of component commonality in a product family, which was later used in an integrated platform in Ref. [45] to support product family redesign. While many of the variety metrics discussed above have shown promising results in different domains, there is a lack of methods to combine metrics sharing common components or to find a unifying formulation which connects variety metrics with theoretical concepts like entropy of a system. There is also a lack of criteria which a new variety metric should satisfy. This paper tries to address some of these gaps.

Measuring Variety in Other Domains

Variety metrics are used in different domains under different names. They are often referred to as diversity metrics, while terms like coverage, breadth or heterogeneity may also be used in different domains. Researchers in domains outside of Engineering Design have measured the breadth of ideation using metrics like the mean pairwise distance between ideas [46] or by manually sub-grouping functions into categories [47]. Over the last twenty years, economists have also become increasingly interested in understanding whether diversity among multiple distinct population groups enhances or impedes a society's economic and social development. To quantify the economic impact of diversity, they also needed to create an index that captures how a society divides into various factions or parts.

Starting from the Gini index [48], economists have used various diversity indices to evaluate the degree of social, economic, cultural, and other dissimilarities among people, regions, and countries. The Gini index was re-interpreted by Simpson [49] as the inverse Hirschman–Herfindahl index. A variety of other statistical metrics of diversity including Shannon entropy [50], effective numbers of species (aka Hills metric), Tsallis number, *etc.* are also commonly used in many fields including information theory (to measure the amount of information conveyed) and ecology (to measure diversity of species). The below paragraphs discuss three of these measurement methods — Shannon entropy [50], Richness [51] and HHI [3].

The most commonly used diversity metric is called Shannon Entropy. Shannon entropy quantifies the uncertainty in predicting the group identity of an individual item that is taken at random from the dataset. Shannon entropy becomes zero when there is exactly one group, that is there is no uncertainty in predicting the type of the next randomly chosen item.

Richness quantifies how many different types the dataset of interest contains and is a popular diversity index in ecology. Although widely used, richness does not take into account the abundances of each type within their group. This can be understood by an example. Suppose we take the same number of species and change the distribution of animals in them. One group has three tigers, three wolves, and three Asian elephants. The second group has seven tigers, one wolf, and one Asian elephant. The first group, with more even number of animals of each species, should get a higher diversity score, however, the ‘Richness’ metric gives

the same score of three to both groups. On careful observation, one may notice that this same problem occurs if one is using design metrics inspired by the ‘Richness’ metric, which count the number of bins in a set of designs [16]. This property (called evenness sensitivity in [52]) is satisfied by other metrics like Shannon entropy, which consider the abundances in each category as well as the number of categories.

The Herfindahl–Hirschman Index (HHI) is a statistical measure of concentration [3,53]. HHI is used by the Department of Justice and the Federal Reserve in the analysis of competitive effects of mergers. It accounts for the number of firms in a market, as well as their concentration, by incorporating the relative size (that is, market share) of all firms in a market. For a market with N firms, HHI is calculated by squaring the market share (MS_i) of all firms ($i \in \{1, \dots, N\}$) in a market and then summing the squares, as follows:

$$HHI = \sum_{i=1}^N (MS_i)^2 \quad (2)$$

Markets with more concentration (less variety) will have a few large square terms. HHI has also been used in other domains ranging from the measurement of linguistic diversity [54] to the measurement of academic specialization [55].

This paper proposes a variant of HHI metric named Herfindahl–Hirschman Index for Design (HHID). The metric is also inspired by Fuge *et al.* [22], who argued that variety metrics are coverage functions which should belong to this family of functions. They introduced a probabilistic model that computes a family of repeatable variety metrics trained on expert data. In this work, our proposed metric also satisfies the properties of submodularity, which allows a practitioner to optimize variety using a greedy heuristic algorithm. The metric does not necessitate finding hierarchical trees, simplifying the variety calculation. The results show that unlike past metrics, this new metric has better alignment with the judgment of variety by people.

While studying HHID, one may ask: why should someone use HHID and not another metric which is a slight variant of it (say cubic power instead of square terms)? To answer this, we show below that the HHID metric is just one instance of a more generalized class of Sharma-Mittal entropy metrics, which can also be used to measure design variety. Commonly used Hartley, Shannon and Quadratic entropy, and the families of Tsallis, Renyi, and Arimoto entropies, can all be derived as special cases of Sharma-Mittal entropy metric [56]. This insight helps unify past notions of design variety under a common mathematical form.

Unifying the Space of Variety Metrics

Sharma-Mittal entropy (SME) is a generalized class of entropy measurement methods that unifies multiple past proposals to measure diversity. It argues that the uncertainty in a discrete random variable $K = k_1, k_2, \dots, k_n$ can be measured by its entropy. Sharma-Mittal entropy (SME) can be defined

as:

$$SME(K) = \frac{1}{t-1} \left[1 - \left(\sum_{i=1}^n P(k_i)^r \right)^{\frac{t-1}{r-1}} \right] \quad (3)$$

where r is the order and t the degree of the entropy measure. The order r is any positive real-valued number except 1. The degree t can be any real-valued number except 1. $P(k_i)$ is the proportion of ideas of variable k . Fig. 1 shows how metrics like Shannon, Quadratic (HHI), Tsallis, Effective number, and others can be obtained using different values of order and degree parameters³.

Although the Equation 3 may not immediately appear intuitive, there are many ways to build understanding of this space of metrics. For example, all of the SME metrics can be thought of as quantifying the average surprise that would be experienced if the value of the random variable K were learned. The order parameter r determines what kind of averaging function is used. r can be thought of as an index of the imbalance of the entropy function, which indicates how much the entropy measure discounts minor (low probability) hypotheses. For example, when $r = 0$, entropy becomes an increasing function of the mere number of the available options. When r goes to infinity, on the other hand, entropy becomes a (decreasing) function of the probability of a single most likely hypothesis. The degree parameter t governs which kind of surprise is averaged. It can be considered as a deformation parameter of the probability distribution [4] and unlike r , it does not have an intuitive explanation. While these relationships between r and t may, at first, appear to just be mathematical curiosities, we show below that by viewing variety in this way, researchers can better scientifically study and uncover how people make decisions about variety—for example, by determining ranges of r and t that agree well with expert opinion.

METHODOLOGY

In this section, we first describe variety measurement methods using the Sharma-Mittal Entropy and then show how a Herfindahl–Hirschman Index based metric can be derived from it. Next, we show an example of variety calculation using the new metric. We show that the new metric can be optimized using a simple greedy algorithm to find sets of ideas with the highest variety. We finally show example computation of variety using the Sharma-Mittal entropy, which generalizes HHI.

The Sharma-Mittal Entropy for Design

In this section, we propose a variant of SME that can measure the variety of a set of designs. To do so, we assume that we are given a set of designs S . As commonly

³Note that limits, which exist, are used for points where the above equation is undefined.

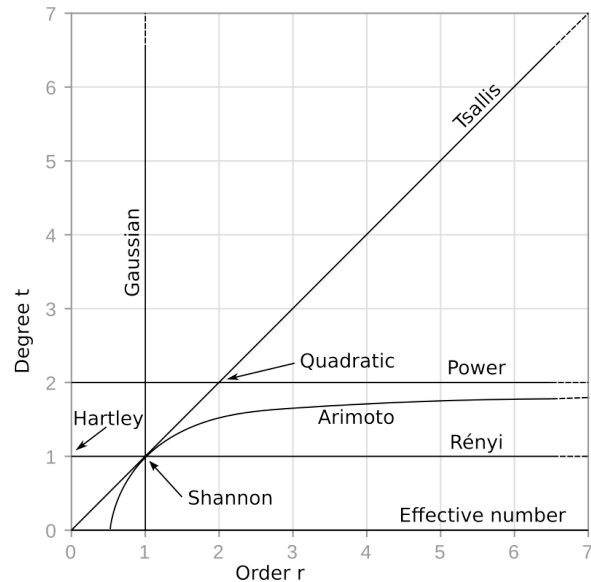


Fig. 1. The two-parameter Sharma-Mittal entropies. Different existing entropy metrics like Shannon, Hartley, Tsallis, Quadratic *etc.* are incorporated in this class of entropies.

used in literature, it is assumed that a design is the represented by a certain level of abstraction like the physical principle, the working principle, the embodiment and the detail level. As explained in Ref. [41], a generated concept of a motor could for instance exist of the ideas “electromagnetism” at the physical principle level, “coils for attracting and repelling permanent magnets” at a working principle level, a schematic or description of the placement of the coils and permanent magnets on the shaft and casing at the embodiment level, and a detailed drawing or description of the parts and assembly at the detail level.

Each design within a set S can be described by a list of attributes (the attributes can be hierarchical levels like functional principle, working principle, embodiment, and detail similar to SVS and NM above or they can be non-hierarchical categorical attributes). We define Sharma-Mittal Entropy for Design (SMED) for each attribute by replacing $P(k_i)$ in Eq. 3 by the corresponding proportion of functional principle. Hence, $SMED_F(S)$ for functional principle is defined as:

$$SMED_F(S) = \frac{1}{t-1} \left[1 - \left(\sum_{i=1}^{N_f} \left(\frac{|FP_i|}{N} \right)^r \right)^{\frac{t-1}{r-1}} \right] \quad (4)$$

Here, $|FP_i|$ is the number of designs using functional principle i and N_f is the total number of functional principles (or the number of categories based on any factor, as defined by a designer). N is the total number of designs in the set S . Similarly, the paper defines $SMED_W$ for working principle, $SMED_E$ for embodiment and $SMED_D$ for details (or any number of attributes defined for a design). The total variety score for a set S is defined as the weighted sum of variety score for each type of attribute as follows:

$$SMED(S) = w_1 SMED_F(S) + w_2 SMED_W(S) + w_3 SMED_E(S) + w_4 SMED_D(S) \quad (5)$$

Here, $SMED(S)$ is the total variety score for a set of designs S . The weights w_1 , w_2 , w_3 and w_4 are used to give different importance to variety of different attribute types and can be set such that the resultant value is always bounded to be less than one (say, by setting the sum of weights to 1). For instance, if all factors are equally important, then one can set $w_1 = w_2 = w_3 = w_4 = 1/4$. This paper assumes here that the total variety of a set is a weighted linear sum of the variety of different attributes found in that set. As discussed later, this assumption aligns well with human judgments and also allows the resultant metric to remain submodular for particular cases. By varying the parameters of SMED, one can measure different types of variety. For instance, if one selects $r = t = 0$, the metric reduces to $SMED_F(S) = \text{Number of unique attributes} - 1$, where the attributes or categories can be the functional principles or subjectively defined categories by an expert. This reduction gives a metric, which is similar to the metric proposed by Linsey *et al.* [16], which counts the proportion of unique bins (categories) to the total number of bins. The next section shows how HHI defined in Eq. 2 is a special case of SMED metric defined in Eq. 4 by using $r = t = 2$. These specific values of r and t are selected as they are the smallest integral values satisfying the optimizability criteria. HHI is a common measure used in domains like economics, and as shown in prior work [57], it aligns well with human interpretation of variety.

The Herfindahl–Hirschman Index for Design The Herfindahl index (also known as Herfindahl–Hirschman Index, HHI) measures a firm’s size relative to the industry and indicates the amount of competition among firms. The mathematical structure of HHI was provided in Eqn. 2. The value of HHI measures the probability that two randomly chosen individuals in a society belong to the same groups⁴. This section proposes a variant of HHI, named HHID, that can measure the variety of a set of designs, where each design is described by a set of attributes. We calculate the HHID index for each attribute separately for the entire set. For example, the HHID index for the ‘functional principle’ attribute type is given by:

$$HHID_F(S) = 1 - \frac{\sum_{i=1}^{N_f} |FP_i|^2}{N^2} \quad (6)$$

One may notice that $HHID_F(S)$ can be derived from $SMED_F(S)$ in Eq. 4 by setting both r and t parameters to two, showing that HHID is a special case of the broader Sharma-Mittal class of metrics. When $N_f \geq N$, $HHID_F(S)$ varies

⁴The interpretation of HHI as the probability that two individuals selected at random from a set represent the same group assumes that the first person is replaced to the set before taking the second person.

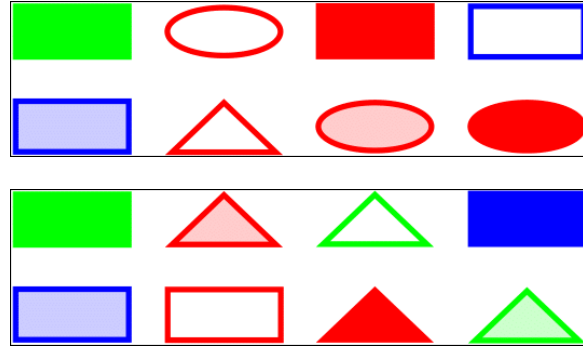


Fig. 2. Example of two polygon sets (Top shows Set A and bottom shows Set B) shown to participants in our experiment. Participant answers the question: “Which set is more diverse?”

between 0 to $1 - 1/N$. Unlike the HHI definition in Eq. 2, Eq. 6 subtracts the value from 1. This definition is closer to the Gini-Simpson index, which is also known in ecology as the probability of interspecific encounter (PIE) [58]. $HHID_F(S)$ ’s value is maximum when all ideas have unique functional principles in the set. Mathematically, it measures the probability that two randomly chosen ideas in the set have different functional principles. Similar to $SMED$, the total variety score for a set S can be defined as the weighted sum of variety score for each type of attribute as follows:

$$HHID(S) = w_1 HHID_F(S) + w_2 HHID_W(S) + w_3 HHID_E(S) + w_4 HHID_D(S) \quad (7)$$

Example variety calculation using proposed metrics

To demonstrate HHID calculation, an illustrative example is discussed next, which is shown in Fig. 2 with two sets of items. In this example, we use polygons instead of a case study from Engineering Design due to two reasons — attributes like shape and color are easy to visualize and one can do one to one mapping of a polygon attributes to the attributes of any Engineering Design idea with a similar number of total attributes.

In Fig. 2, for the set shown on top, there are eight polygons ($N = 8$). There are four items with a rectangular shape, three items with an oval shape and one triangular shaped. There are five red-colored polygons, two blue and one green. Three items have a solid fill, two have shaded and three are empty inside. Without loss of generality, for this example, we assume that color is the functional principle of a polygon, shape is the working principle and shading is the embodiment. It is also assumed that all three levels are equally important in deciding the variety of Set A ($w_1 = w_2 = w_3 = \frac{1}{3}$) and $N_f = 3$ as there are three unique functional principles (color). The $HHID_F$ score for color will be $1 - ((5/8)^2 + (2/8)^2 + (1/8)^2) = 0.531$. Similarly, $HHID_W$ score for shape will be $1 - ((4/8)^2 + (3/8)^2 + (1/8)^2) = 0.593$ and $HHID_E$ score for fill will be $1 - ((3/8)^2 + (2/8)^2 + (3/8)^2) = 0.656$. As all features are equally important, the total $HHID$ for the set of

designs will be $(0.531 + 0.593 + 0.656)/3 = 0.593$. Similarly, the variety of any set of designs can be calculated. For instance, the $HHID(S)$ score for the set at bottom is $(0.656 + 0.500 + 0.656)/3 = 0.604$, using Eq. 7. These two sets are close in their variety scores using $HHID$, with the bottom set having a slightly higher score than the top set. While the top set lacks in color variety (five red polygons), the bottom set lacks in shape variety (no oval shape).

Next, the variety score of both the sets using $SMED$ is calculated. For demonstration, we use two settings of order and degree parameters. First, both values are set to zero, *i.e.* $r = t = 0$. In this case, the $SMED_F$, $SMED_W$ and $SMED_E$ are each one minus the number of unique principles (three). Hence, the $SMED(S)$ is six for the top set. The bottom sets $SMED(S)$ is five (as it lacks one type of shape). This shows that the top set has a higher variety score, if $SMED$ parameters are set such that they give more importance only to the unique number of groups found. The second setting we show is $r = 4$ and $t = 2$. The $SMED(S)$ scores for the top and bottom set are 0.558 and 0.598 respectively, again giving a higher variety score to the bottom set, similar to $HHID$. In all these cases, it is assumed that variety for color, shape, and shading are equally important. However, suppose a problem requires that the shape variety is four times more important than color and fill, then we can set $w_1 = w_3 = \frac{1}{6}$, $w_2 = \frac{4}{6}$. In that case, we get variety scores of 0.561 and 0.549 for the top and the bottom set respectively. The first set, which has three unique shapes, gets a higher variety score if more importance is given to the variety of shape. Thus by changing the weights for different attributes, one can customize the variety metric to meet the demands of a particular domain.

Optimizing variety of a set

Using metrics like SVS, NM and $HHID$, one can measure the variety of a given set of ideas (like the sets shown in Fig. 2). However, what happens when one wants to choose a small set of polygons (say five) which have the maximum variety out of a thousand items? One way is to enumerate all possible sets of size five (more than 8 trillion sets for a ground set of 1000 items), calculate their variety score and then find the set with the highest variety score. This approach becomes intractable as the number of items in the ground set increases.

Another approach, and the one used in this paper, is to leverage mathematical properties of the variety function and find approximate solutions close to the optimal. This $HHID$ metric is a submodular set function. Submodular functions are functions defined over sets that are designed to model diminishing marginal utility, which is the mathematical property one needs to model diversity or variety [22]. Having the submodularity property means that the variety metric follows the law of diminishing returns — when a design is added to a larger set, the increase in $HHID$ score is smaller compared to the case when the same design is added to a smaller set. This property can be exploited to find sets of maximum variety using a greedy algorithm [59], which guarantees that the variety of the greedy search solution will be within 63.2%

(or $1 - \frac{1}{e}$) of the variety of the optimal solution.

To find sets of maximum variety, one can use a submodular greedy algorithm explained in [59]. Given the set V of all ideas, the algorithm starts with an empty set $S = \{\}$ and add ideas to this set which give the maximum marginal gain in the submodular function. At every step, it adds one idea at a time, such that the selected idea $i \in V$ is the one with the highest marginal gain $\delta HHID(S \cup i)$ on set S . At each step, the algorithm adds the idea that will give a maximum increase in variety in the set S . Finally, as the function in Eq. 7 is submodular and monotonic, the algorithm is also theoretically guaranteed to provide the best possible $(1 - \frac{1}{e})$ polynomial-time approximation to the optimal solution [60, 61]. $SMED_F(S)$ function defined in Equation 5 is concave as long as $t \geq 2 - 1/r$ (see [62] for a proof). This implies, in particular, the concavity of all metrics lying on the diagonal line in Fig. 1, which are obtained by positing $r = t$ is also concave. Metrics of this kind are often labeled after Tsallis's [63] work in generalized thermodynamics, which has also been advocated as a compelling approach to the measurement of biological diversity [64]. As a sum of concave functions over modular functions is submodular [65], the resultant $SMED(S)$ metric is also submodular for $t \geq 2 - 1/r$. Hence, the key takeaway is that one can optimize any SME derived metrics for all values $t \geq 2 - 1/r$ in polynomial time using a simple greedy algorithm. This includes $HHID$, which is obtained by setting $r = t = 2$.

EXPERIMENTS AND RESULTS

We conducted an experiment to benchmark the proposed $HHID$ metric with the commonly used SVS and NM metrics using a known and easily verifiable ground truth based on polygons. Next, another experiment is reported, which uses milk frother design sketches provided by engineering students and rated by domain experts. Before introducing our experiment and its main results and implications, we describe how the experimental dataset of set comparisons was constructed. As shown later, constructing such sets is non-trivial, and one contribution of this paper lies in describing a procedure for constructing such comparison sets for new domains.

Estimating Design Variety Ground Truth using Human Pairwise Comparisons

The first step in vetting design rating metrics is to identify a 'ground truth' of the measure that the metric is trying to capture and then calculate how accurate any given metric is in capturing that ground truth. However, for the case study presented here (design variety), ground truth estimation is difficult due to the large combinatorial space for sets of items and the lack of a benchmark dataset. For instance, a small set of thirty design ideas has more than one billion possible sets of designs for which variety needs to be calculated. Exhaustively calculating the ground truth for all designs is infeasible. To avoid circularity, any existing variety metrics are not used to create the ground truth. Doing so would assume

that a given metric represents true variety, which is what the ground truth is used to establish. Instead, this paper proposes the development of a ground truth by directly asking human raters.

To establish a ground truth dataset for calculating the design variety, three components are needed:

1. A ground set of design items over which sets are created
2. Sets of designs derived from the ground set for which variety scores are calculated
3. Tree annotations for each design item to enable the calculation of tree-based metrics

Variety scores are calculated on a set of designs. However, human raters are not good at giving absolute scores [66] due to differences between internal scales of subjects, a well-known problem for subjective scoring. For instance, given the set of designs shown in Fig. 5 top, it would be difficult for a human rater to say whether this set of six designs scores 6 out of 10 or 8 out of 10 for variety. Different raters may also use different internal scales.

In contrast, if a rater is asked to rate whether they find the variety of set shown in Fig. 5 top set greater than the variety of those shown in Fig. 5 bottom set, they may answer it relatively easily because humans are better at comparing items than giving absolute scores [67]. Hence, this paper proposes that a ground truth dataset for variety should be created using pairwise queries (ordinal judgments), where each query contains two sets and there is a consensus among human raters that one set has higher variety compared to the other set. To elicit responses from experts, two sets at a time are given to them and they are asked pairwise comparisons of the form: “Which set of designs has higher variety?”

Measuring Variety for Polygons

In this experiment, the performance of SVS and NM metrics in measuring the variety of a set of polygons is compared with HHID, which is a special case of Sharma-Mittal entropy. Initially, a base set of 27 polygons is created. Each polygon has three attributes — shape, color, and shading. Each attribute can take three unique values. Polygons can be rectangular, triangular or oval-shaped. They can be red, blue or green colored. Shading varies between polygons as complete fill, shaded or empty.

The polygon example, which does not represent a real-world design, is intentionally chosen to compare design metrics. The difficulty with using a real-world example to establish a ground truth for quantitative metrics is that such examples have many moving parts and human judges have low-agreement on what attributes should be extracted from the design and which ones are important in determining their similarity. For real-world examples, due to the inherent complexity in the measurement of attributes and design performance, it is difficult to say conclusively whether the lack of alignment of human judgment with a variety metric is due to the wrong choice of attributes or the wrong choice of the method measuring variety over those attributes. While the polygons example does not represent an actual engineering

design solution, it is used to compare metrics when there is no ambiguity in design attributes (shape, color, fill). Our argument is that metrics that can measure variety for many complex domains should at least fair well in measuring variety for a simpler polygon-based ground truth dataset. Later sections provide a more complex example and discuss the issue with capturing attributes.

The total number of possible sets of polygons is large (2^{27}), hence calculating the variety score of all possible sets is time-consuming. Instead, the search is narrowed down to focus on three set sizes: when the number of items in a set is four, six and eight. The researchers observed in their preliminary experiments that if human raters are asked to compare sets with larger than eight items, the task becomes too difficult for them, as evident by low agreement between different raters. For a given set size (say size six), the total number of ways two sets can be compared is also quite large (more than 43 billion set comparisons). Hence, we first randomly select 100 sets for comparison. From these 100 sets, we calculate all possible pairwise comparisons (4950 comparisons with each comparison containing two sets of size six). Next, we calculate SVS, NM, and HHID scores for all the sets in each comparison. For SVS and NM computation, the analysis assumes that ‘Color’ is the functional principle, ‘Shape’ is the working principle and ‘Shading’ is the embodiment.

Result 1: Existing metrics cannot distinguish between sets. Table 1 shows the percentage of comparisons where each metric finds both the sets of equal variety. Note that SVS and NM metrics do not distinguish between a large percentage of comparisons (31.7% and 21.4% for sets of size six), while HHID gives identical scores to a much smaller percentage of pairwise comparisons (14.7% for sets of size six). This implies that existing metrics are not sensitive or discriminative to differences between sets.

Result 2: Existing metrics vote similarly to one another. Table 2 shows the percentage agreement between different metrics. SVS and NM vote similarly for 80-85% of set comparisons for various set sizes. This means that for a large proportion of comparisons, both metrics are indistinguishable as they give the same pairwise response. If SVS finds Set A has higher variety, then so does NM. In contrast, the agreement between HHID and other metrics is close to random. Due to the lack of a benchmark dataset, it is difficult to comment on whether a lack of agreement between metrics is a good thing or not. We show later in the results that HHID aligns with the human raters more than SVS and NM.

Establish the ground truth for comparing different metrics required the following steps. First, pairwise comparisons where SVS and NM could distinguish between the two sets were selected; that is, both the metrics did not calculate the same variety score for both sets. This is important since we want any collected human judgment to differentiate existing metrics, and we cannot do this if we select comparisons where the two metrics calculate the same value. Secondly, the sets where both metrics disagreed on their vote are selected. This means if SVS voted Set A to be higher variety,

Method	Same Score		
	SVS	NM	HHID
Size 4	27.3%	37.0%	15.8%
Size 6	31.7%	21.4%	14.7%
Size 8	28.5%	12.9%	10.9%
Size 10	31.2%	14.5%	9.2%

Table 1. Percentage of pairwise comparisons when design metrics give same score to both designs. Lower percentages are good as it indicates that a metric can distinguish between sets. SVS metric gives same score for approximately 30% of the sets.

Method	Agreement		
	SVS-NM	HHI-SVS	HHI-NM
Size 4	84.4%	54.2%	50.2%
Size 6	81.0%	47.6%	50.0%
Size 8	82.5%	49.4%	56.9%
Size 10	84.4%	54.2%	50.2%

Table 2. Agreement between metrics for pairwise comparisons. SVS and NM tend to vote similarly for more than 80% of the sets.

then NM would give Set B a higher variety score. Note that this is a small set of pairwise comparisons — as we noted from Table 1, both metrics vote similarly for more than 80% of the comparisons and tend to give same scores to a large percentage (up to 37%) of the sets.

Finally, the top 5 sets where SVS is most confident that one set has higher variety than another are selected and the top 5 sets where NM is most confident that one set has higher variety than another set (*i.e.*, the difference between the scores are maximum). We combine these two to generate 10 queries which are then given to human raters. Finding human annotations for such sets allows a researcher to find out which of the two metrics better aligns with human responses.

To find the ground truth for polygons, an Amazon Turk study was conducted to collect responses from crowd workers for pairwise queries. A sample query with two sets of eight polygons is shown in Fig. 2. Judging the variety of polygons does not require expertise in the area and Amazon Turk enables getting a large number of responses quickly. We collected pairwise responses for three different set sizes. For each set size, ten pairwise queries were created. For each query, ten responses from Amazon Turk participants were collected. To ensure the quality of responses from the crowdworkers, the following steps were taken: 1) The order of the queries was randomized and also the order of the options shown to different participants to reduce the possibility of any ordering bias. 2) The surveys were divided into two parts to reduce fatigue. 3) No worker was repeated across surveys, and 4) Six queries were repeated to filter out workers with very low internal consistency.

Result 3: Human raters largely agree on what it means to have a high variety set of polygons. The survey responses showed that on average people had consensus on one set being more diverse or higher variety than another set. The number of votes received by the set pairwise query receiving a majority vote for sets of size four was: [9, 8, 9, 7, 6, 9, 8, 6,

8, 7] respectively. This means that for the first query, 9 people out of 10 voted for the same set. For the second query where two sets of size four were shown, 8 people voted for the same set as being of higher variety. Similarly, for sets of size six, [5, 5, 9, 9, 8, 6, 8, 5, 8] votes were received by the majority set and [7, 5, 7, 7, 9, 9, 8, 6, 7, 6] votes were received by the majority set for sets of size 8.

A direct comparison between SVS, NM, and HHID metrics using the published weights would be unfair to SVS and NM, as HHID weight parameters can be optimized specifically for each domain. The published weights for SVS metric is [10, 6, 3, 1] and published weights for NM metric is [10, 5, 2, 1]. To maximize their performance, SVS and NM metrics are given the same flexibility by allowing the weights of functional principle, working principle and embodiment to be optimized. For a given metric (say SVS) and weight combination (say 4, 3, 3), the variety scores for both sets in a given pairwise comparison are calculated. Suppose there are ten humans who voted on a pairwise comparison task. If SVS metric finds that Set A has more variety than Set B, and eight humans had also voted this way, then eight points are allocated to the SVS metric. If the metric found Set B has higher variety than Set A, then this metric receives the two points which humans gave to the other set. As we ask 30 different queries from people, to judge the metric, the aggregated points for all 30 queries are calculated for each metric.

Based on how people voted in this experiment, the maximum number of points that any metric can receive is 220 — that is if it always votes with the majority opinion of human raters. Now, suppose a metric receives 200 points in total, then we say that it has 90.9% alignment ($100 \times 200 / 220 = 90.9$) with human ratings.

Result 4: HHID outperforms SVS and NM w.r.t. human agreement on polygon variety. Table 3 shows the comparison between SVS, NM, and HHID for alignment with human ratings. SVS and HHID have similar best-case performance for this dataset. By varying the weights of each functional level between 1 to 10 in steps of 1, gives a 1000 possible performance scores corresponding to each weight combination [w1, w2, w3] (this is in contrast to using a fixed combination of weight, like [10, 6, 3, 1] for SVS). The results show that HHID performs better than SVS in the median case, where the median is calculated over all the thousand weight combinations.

Table 3 shows that HHID aligns with human perception of variety to the highest degree, irrespective of the choice of weights — that is, its performance is robust to weight choices. Even in the worst case, HHID aligns with 74.5% of human ratings. We find that the highest performance is obtained for many combinations of weights like 1, 2 and 10. On first glance, SVS also seems to perform well for the median case. However, this does not mean SVS is suitable to measure variety as we only select the queries where SVS is able to differentiate between the two sides. It is also important to recall that the comparisons were generated such that SVS has high confidence in its choice between both the sets (by design). In contrast, if we select sets to compare at random,

SVS calculates the same score for more than one-fourth of the queries. This drastically reduces the SVS performance in alignment with human responses — humans generally showed a clear preference between the variety of two sets, but SVS would be indifferent. Hence, due to the better accuracy and higher sensitivity, the HHID metric outperforms both SVS and NM in alignment with human’s judgment of variety.

Method	Median Case	Best Case	Worst Case	Sample optimal weights
HHID	81.8%	95.4%	74.5%	1, 2, 10
SVS	79.0%	95.4%	59.0%	2, 1, 1
NM	54.5%	86.3%	40.9%	10, 3, 1

Table 3. Comparison of design variety metrics in alignment with human ratings

Result 5: Sharma-Mittal Entropy parameters show most entropy metrics align well with human judgements This section shows how SME aligns with human perception of variety. To do so, the weights of the metric for color, shape, and shading are set to be equal and only the order r and degree t of the SME metric are varied.

r and t are varied between 0 to 10 at steps of 1 (note that r and t are not necessarily restricted to integral values). The results are shown in Fig. 3. For each combination, the resultant alignment with humans is calculated. As shown by the white region, the metric achieves the highest performance for multiple combinations of r and t , including $r = t = 2$ used to define HHID. This leads to the question: What does this indicate about how people think about measuring variety of polygons?

To dive into this question, it is important to first understand the parameters of the SME metric. In the SME metric, the order parameter r is an index of the insensitivity to less abundant principles. As r increases, variety gets closer and closer to a simple (decreasing) function of one single element in the distribution, which is the relative abundance of the most common principle (most common color, shape or shading in the set). When $r = 0$, on the contrary, variety becomes an increasing function of the plain number of principles with non-null relative abundance (*e.g.*, the count of the number of unique shapes, colors or shading). This shows that the order parameter r indicates how much a variety measure disregards relatively rare principles. The higher r is, the more the common categories are regarded and the rare categories are discounted in the measurement of diversity. The role of the degree parameter t is more technical: it affects a few important metric properties, which is elaborated in detail in literature on mathematical analysis of SME [4].

The results show that the metric top performance is indifferent to variations in r . This means some people may have focused on just the count of classes, while others may have focused on the largest class. Performance is sensitive to values in t , with a decrease in performance as t goes above two.

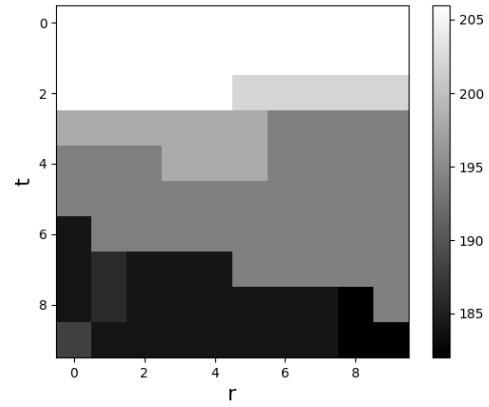


Fig. 3. Plot of performance for different values of order (r) and degree (t) parameters of Sharma-Mittal Entropy. Performance is high for many common entropy metrics like Shannon entropy ($r=1$, $t=1$) and HHID ($r=2$, $t=2$).

Result 6: We can find sets of designs with highest variety

One of the auxiliary outcomes of using an HHID derived index for variety measurement is that it provides a simple method to find the highest variety sets. Suppose in an ideation exercise, 10 teams get together to generate a total of 27 ideas. Our goal is to combine all ideas into one large set and then down-select to a small set of ideas that provide a distribution over the design space. To demonstrate the concept, we assume that ideas produced by all the teams are represented by the 27 polygons discussed before and the goal is to find a subset of five ideas, which have the highest variety (one can pick any size of the subset).

If one wants to find the subset of size five ideas with the highest SVS variety (or any other tree-based variety metric), they will have to calculate all possible combinations of five ideas, then calculate the tree for each subset, estimate SVS scores for each set and finally pick the set with the highest SVS score. This exercise will require enumerating all 80,730 (27 choose 5) possible trees for each set of five polygons. This approach becomes infeasible when the ground set becomes large due to a large number of possible options (mathematically, this is because the problem is NP-Hard).

In contrast, we use a greedy algorithm [59, 68] to rank order all polygons or to select a subset. For size 5, it will require only 125 evaluations, which requires 99.84% fewer calculations. When applied to polygons, the resultant set, with highest variety for color, shape, and shading, is shown in Fig. 4. The method selects one polygon at a time based on which polygon provides the highest marginal gain. As mentioned above, this is possible in polynomial time due to the submodular behavior of HHID. A practitioner may also wonder how many ideas should provide sufficient coverage over the design space. While this work does not show how many ideas are enough to explore the design space, past work [68, 69] has shown straightforward methods to estimate the cutoff for the number of ideas needed using the marginal

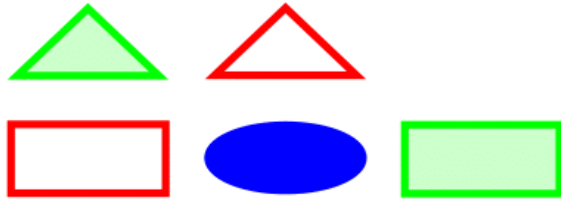


Fig. 4 Set of five polygons with highest variety found using a greedy algorithm applied to submodular objective function

gain of a submodular function. The same method applies to the current variety metric due to its submodularity and can be used to find the size of the subset.

Measuring Variety for Milk-Frother Sketches

While the polygon example discussed in the previous section helped to validate the metrics, using a problem with little complexity, it did not represent an actual engineering design solution. In this section, an additional experiment of an engineering design problem is discussed, where the goal is to measure the variety of early-concept design sketches. We use experts to judge items from a pre-existing dataset of milk-frother sketches to create a ground truth dataset comprising of pairwise comparisons. Finally, we measure how well different variety metrics align with the ground truth to measure their accuracy.

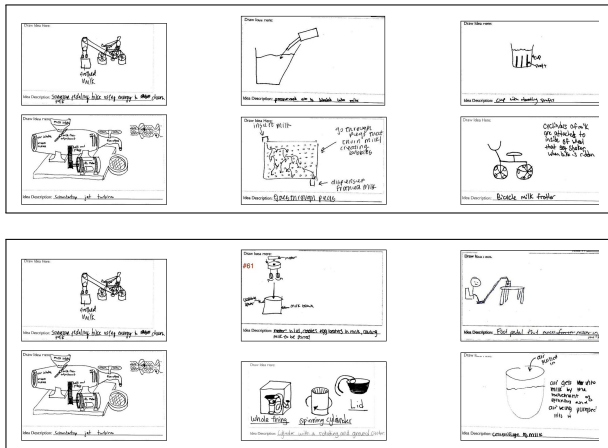


Fig. 5. Top: Sample of Set A where all raters agreed it was more diverse than Set B. Bottom: Sample of Set B where all raters agreed it was less diverse than Set A.

To measure the variety of milk-frothers, data from a previous experiment conducted by Starkey *et al.* [70] is taken, which consisted of 934 idea sketches. Specifically, the data set consisted of ideas developed by 89 first-year students from an undergraduate engineering course and 52 senior students from a capstone engineering course including 95 males and 46 females. The ideas developed in this dataset were from a design task where participants were asked to generate

ideas for a “novel and efficient milk frother”. This task was selected because the task addressed solving a product-based problem.

In order to calculate the metrics based on hierarchical features, the results from the previously developed Design Rating Survey (DRS) was used to classify the features addressed by each design concept (see [71] for more details). Twenty questions on the DRS were used to help raters classify the features each design concept addressed. The results of the DRS were then split into which category they addressed in the extension metrics: physical principle, working principle, or embodiment. The physical principle was determined by what type of power source was used for to power the product (i.e. manual, battery). On the other hand, the working principle was determined by what type of motion was used by the product (i.e. stirring, shaking) and the embodiment was determined by what the product looked like (i.e. shake weight, handheld frother).

For this study, to create the dataset of sets of milk-frother sketches, a ground set of ten design sketches is adopted from Ahmed *et al.* [72]. The benefit of using these ten sketches was the availability of hierarchical features as well as information in the form of subjective idea maps, which is later used for gaining additional insights. The total number of possible sets which can be formed using these ten sketches is $1024(2^{10} \text{sets})$. Similar to the polygon case, the goal is to find a small set of pairwise comparisons of sets, which humans agree on. It is important to create a ground-truth dataset of pairwise queries where human input is most useful in distinguishing between well-known metrics. The process of identifying pairwise comparisons which are shown to human experts is described next.

From the ten sketches, pairwise queries with sets of six sketches have to be created. We decided to create the ground truth with pairs of six sketches as the median number of sketches made by a participant in the entire milk-frother dataset [70] was six. The number of unique sets of size six is 210 (10 choose 6). To see the distribution of variety scores and guide the selection of a ground truth dataset, the variety scores for all these sets using SVS and NM metric is calculated. However, in this case, the information about Euclidean embeddings for each sketch as discussed in [72] is also available, which is used to guide the selection of queries. These embeddings are essentially 2-D maps with each design having x and y coordinates allocated to them. Similar designs occur closer to each other than dissimilar designs on this map. To decide which sets of six sketches to ask humans to rate, information from three metrics (SVS, NM and average pairwise distance) is used. The last metric is derived using an embedding of designs derived in the study by Ahmed *et al.* [72]. One design embedding was picked randomly (as each participant in the study had a different design embedding and only one design embedding was needed to guide our experiment) and it provides the 2-D positions for each sketch. The choice of the design embedding does not alter the key findings of this section as it is only used to guide the selection of queries which are asked from people. The variety scores for all 210 sets was calculated and all sets

were rank-ordered from the highest variety set to the lowest variety set, where the variety is measured using the pairwise average distance metric.

Out of these 210 sets, we obtained 21,945 pairs of sets (210 choose 2) and calculated the absolute rank difference between the two items for each pair. A small rank difference implies that the two sets in a comparison have similar variety, while a large rank difference implies that the metric is confident that one of the set has a significantly higher variety than the other. After calculating the rank differences, 20 comparisons were selected based on two factors. First, comparisons were selected where each metric (SVS, and NM) votes differently on which set has higher variety — *i.e.*, if all ratings agree on the comparison, then human expert ratings would not discriminate them. Second, sets with a high-rank difference, but that also differ from sets we are using in other selected comparisons were selected. This ensures that a metric is confident in its vote, but also the queries provide good coverage over different types of sets in the data by ignoring pairs that have already been selected.

Among these candidate sets, 20 pairwise queries are selected that are given to four expert raters using a Qualtrics survey. Sets from a typical query is shown Fig. 5. Two comparisons (10% repeated queries) were repeated in each survey to measure the internal consistency of each expert, and a total of 22 queries were given to them. Experts can choose whether Set A is higher variety compared to Set B or they can select the option of ‘Can’t decide’. On average, the raters took 24 minutes to complete the survey. From these expert ratings, we find that all four experts agreed on 9 out of 20 queries, while at least three experts agreed on 15 out of 20 queries. Due to a majority agreement among experts, these 15 queries are selected as the ground truth dataset for comparing variety metrics. Next, they are used as ground-truth dataset to compare variety metrics.

Result 7: SVS and NM are equivalent to random chance, w.r.t. matching expert assessments of milk-frother variety. After finding the relative variety scores for each query using SVS and NM, it is seen that they align with only one-third (33.3%) of the ground truth dataset — that is five comparisons. To see if this low performance is due to specific choice of weights, 1000 possible weight combinations for SVS and NM are tested to report how close these metrics are to human experts. To explore the sensitivity of these results, we calculate the NM and SVS scores for every valid weight combination used by each metric. Using these weights, we find that SVS aligns with 33.3% of the pairwise expert assessments of milk-frother variety irrespective of the weights used — that is, changing the tree weights used by SVS has zero effect on whether or not it agrees with human experts. NM aligns with 33.3% of the dataset for 95.6% of all the weight combinations. For the rest, it has no alignment with any expert ratings — that is, NM’s scores are more sensitive to its internal weights, but not in a way that benefits its score accuracy. The alignment scores are close to random chance for three categories (Greater, Smaller and Equal) showing that SVS and NM are unable to capture human intuition of

variety for the examples we tested.

Result 8: HHID robustly outperforms SVS and NM w.r.t. human comparisons, but still has a non-trivial error. In contrast to SVS and NM, HHID aligns with 9 out of 15 comparisons when weights are optimized for each level. We find that many weight configurations for HHID lead to highest performance, including $w=[1, 9, 5]$.

Hence, HHID aligns with the human judgment of variety more than both SVS and NM metrics for two standard datasets. However, it still is not 100% aligned to human benchmarks. However, in the second experiment, we had assumed that the annotations provided for SVS, NM, and HHID for different hierarchical levels are accurate. If this is not the case, any variety metric will have a large error as it may not capture the true features. Constructing the hierarchical trees is outside the scope of this paper but it is important to understand that metrics may be limited by the specific choice of how one constructs a tree, which also needs to be verified.

We propose that by using our above method for constructing these ground truth variety comparisons, future papers will be able to use these and other ground truth variety pairwise comparisons to judge the comparative quality of other metrics as well. This would provide a common scale over which metrics are compared.

DISCUSSION

Above experiments highlight several broader implications, both around how variety metrics are constructed and verified, as well as in how existing metrics are used across domains.

Selecting appropriate validation sets for variety metrics is non-trivial

As we showed above, selecting exactly which sets of designs to show experts for ground truth labeling is non-trivial. First, the combinatorial nature of the problem (sets of designs) makes exhaustive labeling by experts impractical for anything above a handful of designs. But randomly subsampling this combinatorial set does not solve the problem: many metrics may trivially agree on a large portion of the space.

We proposed possible desiderata on what comparisons to show experts, as well as several potential methods to make this selection, such as maximal rank order disagreement, distances over embedded spaces computed via past techniques [72], and space coverage over different sets. Constructing comparisons in this fashion does lead to potential bias: as we saw in Result 4, by preferentially sampling sets where metrics were confident in their answers, we may overestimate their performance.

The trade-off here is one of time and cost. If one picks comparisons to maximize discriminative power among metrics, this will inevitably ignore portions of the space where they agree and inflate performance metrics. In contrast, if

one does not do this one may collect many expensive expert comparisons that, while covering the space well, do not provide much value in separating good metrics from bad ones.

One limitation of our proposed approaches is that we currently provide no theoretical guarantees regarding the number or scope of queries needed to achieve a certain assessment accuracy. The number of comparisons we collected above was driven by primarily practical concerns — how many expert comparisons could we realistically expect to collect in our available time budget? Future work could address how to perform this collection optimally (*e.g.*, using Active Learning) and to bound the number of comparisons one would need to collect.

Good variety metrics need to be accurate and discriminative

As we showed in Results 1 and 2, good metrics need to not only be accurate but also highly discriminative or sensitive. We found that commonly used metrics can lack sensitivity across a broad range of comparisons. Even if such metrics are accurate, they have limited usefulness as measurement instruments — that is, they cannot detect small effect sizes in terms of differences in variety. We argue that, in addition to focusing on accuracy, future metric development should compute and account for the sensitivity of the measurement instrument for the given domain, and such quantities should be reported in subsequent papers.

Metric performance can differ significantly across domains

Comparing Results 4 and Result 7, we see that a given metric applied to one domain/problem may have drastically different performance. In our case, SVS performed well for human comparisons on the polygon case, but poorly on the milk-frother case. While it is perhaps obvious that a metric's accuracy depends on where it is applied, we note that, in practice, past researchers have broadly used existing metrics (both SVS, NM, and others) with limited to no verification and calibration of the measurement instrument to that domain.

We believe that our results here should give other researchers pause before blindly applying an existing variety metric to a new problem without first conducting some of the pairwise verification we detail above. We are releasing both the datasets we collected in this paper and the tools we used to construct human comparisons in the hope that future researchers will have an easier time constructing verification tests for new metrics or domains.⁵ We believe that the proposed metric can be used in combination with other design metrics to provide insights from different perspectives of a set of designs. The usage of this metric and creation of new ground truth datasets should take into account the context that designers have deep knowledge in a field and can judge variety through different lenses and with an experience that may not always be possible from a quantitative metric.

Sharma-Mittal entropy is a promising alternative metric that allows optimization of variety

We demonstrated via Results 4 that using HHID matched or exceeded the performance of commonly used metrics. This was true in both the Polygon and Milk-Frother experiments. Calculating the HHID is computationally simpler to the benchmark tree-based constructions of SVS and NM.

More importantly, the submodular form of HHID allows one to efficiently (*i.e.*, in polynomial time) approximate the highest variety sets of designs, given a corpus. For design corpora larger than approximately 50 designs, this leads to orders-of-magnitude reductions in computational effort in finding optimal variety subsets of design, compared to existing metrics. The fact that HHID can be easily optimized to match human judgments for a domain makes it flexible to apply to different problems if one gathers pairwise comparison data as described above.

We further showed in Result 5 that many different settings of Sharma-Mittal entropies are suitable to measure design variety, HHID being one instance of them. We also discussed the conditions under which SME is also submodular, which helps in the optimization of the metric. Different domains tend to use different metrics. This generalization helps one understand why one metric may be more suitable to a particular domain.

It is important to understand a few major assumptions in using quantitative metrics. First, SMED is defined for categorical variables (like red, blue, and green), where all categories are assumed to be equally distant from each other. This assumption is used in NM and SVS too. However, it is possible that in some applications, items may have attributes which are real-valued or a few categories can be more similar to each other than others. In such cases, SMED will not be a suitable choice and future work will explore extending SMED to continuous domains. Second, finding the right attributes (or design representation) is critical to the success of any quantitative design metric. Many manual and automated methods exist to identify suitable attributes for a set of designs. For example, text-based designs may use keyword extraction or topic modeling to identify attributes. Image-based designs may use image descriptors and CAD models may use shape descriptors for attribute identification. This work assumes that the attributes are provided and estimate variety score for the given set of attributes. Identifying the right attribute to represent different designs is outside the scope of our work.

As SMED and HHID are both derived from entropy metrics, theoretically they can give an absolute score about the variety of a system (in this case, a set of ideas). However, in practice, they are relative metrics. This is because the variety score is dependent on what attributes or categories are considered in the evaluation (We use N_f categories in Eq. 4). If one introduces more categories and reallocates ideas to these new categories, then the variety score may change. For instance, if all 'uninteresting' designs are allocated to the same category (given the same attributes), then they will have a small score. However, if one chooses to allocate them dif-

⁵<https://github.com/IDEALLab/design-variety>

ferent attributes, then the variety score will be large. This limitation, which also exists for other quantitative metrics, is an artifact of finding the right attributes, and not necessarily of how the metric is defined.

In comparing sets relative to each other, this paper also assumes that variety measurements are transitive for a fixed set of attributes (or categories). This assumption is backed by the information-theoretic interpretation of Sharma-Mittal entropy, from which our metric is derived. However, it is possible that human variability judgements are not always transitive and this assumption was not explicitly tested in this paper. As this paper uses consensus on pairwise queries to score metrics and not to rank order all sets, this assumption would not affect our results for metric accuracy.

Our future work will focus on using machine learning methods to identify a set of attributes, which are most important in variety estimation. As human judgments are often expensive, an interesting avenue of work will be to cast the fitting of HHID or SME as an active learning problem and reduce the number of expert comparisons needed to adapt design metrics to a new domain. In future work, we will also verify whether human variety judgments are transitive or not.

Possible applications of Sharma-Mittal entropy beyond sets

Morphological matrices are a powerful tool for generating ideas, based on potential variations in a problem's attributes. For a morphological matrix, the variety score can be calculated in different ways, depending on what the end goal is. The morphological matrix is a simple and powerful tool that enables a design engineer to organize and generate all the different alternatives before identifying the best design solution [73–75]. A possible extension of SMED is to morphological matrices. One option is to calculate the variety of the entire matrix, which will inform us how widely do all solutions explore the design space. Another option is to calculate the variety within each function, by listing all idea combinations in it and optionally clustering them. In future work, we will explore how variety-metrics can be integrated with morphological matrices and compare different ways of doing so.

In applications of morphological matrices and many other design exercises, ideas are often grouped together or chunked during the activity. If ideas are chunked into a set of N_f requirements (or categories/clusters), the SMED score can be calculated using Eq. 4. Hence, the metric allows for chunking of ideas into groups. If more categories to which an idea can belong are added, it effectively means an increase in the value of N_f in Eq. 4. This means, for the same set of ideas, adding new categories will lead to an increase in the variety score of the set (assuming new categories are not empty), while reducing the number of categories will lead to a reduction in the variety score. In the extreme case, when there is only one category, variety score is zero. Finally, it may also be needed that the variety score is calculated for one set of attributes (or requirements) and later, a different set of attributes is used to calculate the variety score. One

can also calculate variety scores using the SMED metric for different subset of attributes using Eq. 4. In Figure 2, one may use only shape as the attribute and after the polygons are colored, may use the color as an attribute to calculate variety of sets. However, it is important to note that two scores calculated using different set of attributes cannot be compared with each other meaningfully as the score magnitude also depends on the total number of attributes chosen.

When ideas are collected using methods including brainstorming, the Gallery method, Storyboarding, *etc.*, it is possible that there are missing or incorrectly reported attributes. This situation cannot be handled by existing quantitative metrics, including SMED. A challenging, albeit important, area of future work is to study design metrics under uncertainty in attribute measurement.

CONCLUSION

In this paper, we contributed: (1) a generalization of design variety metric based on the Sharma-Mittal entropy, for which Hirschman-Herfindahl index for design is a special case; (2) a practical procedure for comparing variety metrics via constructing ground truth datasets from pairwise comparisons by experts; and (3) empirically demonstrating the procedure and metric on two new two ground truth datasets using milk-frother design sketches and polygons. Using this dataset, we then compared the performance of two existing and commonly used tree-based metrics and showed that our newly proposed metric aligns with human ratings more than existing metrics. As an ancillary benefit, we also show that by using a simple greedy algorithm, our new metric can find sets of designs with the highest variety in polynomial time.

Overall, our results shed light on some qualities that good design variety metrics should possess and the non-trivial challenges associated with collecting the data needed to measure those qualities. These results guide how and when various commonly used metrics may or may not be valid, as well as a concrete scientific process by which to gain further insight into when and where metrics apply.

We hope that the procedures we outline here can provide a catalyst for deeper discussion regarding how we measure and verify variety within engineering design. We encourage researchers to build upon and contribute to the datasets we have started collecting and distributing for these problems. We hope that by better understanding how to measure the variety and ultimately optimize variety, we will be able to reliably and scalably support designers in improving their creativity and competitiveness.

ACKNOWLEDGEMENTS

This material is based upon work supported by the National Science Foundation under Grant No. 1728086. We acknowledge the effort of both the MTurk workers and expert raters who help us collect ratings.

References

- [1] Shah, J. J., Smith, S. M., and Vargas-Hernandez, N., 2003. "Metrics for measuring ideation effectiveness". *Design studies*, **24**(2), pp. 111–134.
- [2] Nelson, B. A., Wilson, J. O., Rosen, D., and Yen, J., 2009. "Refined metrics for measuring ideation effectiveness". *Design Studies*, **30**(6), pp. 737–743.
- [3] Hirschman, A. O., 1964. "The paternity of an index". *The American economic review*, **54**(5), pp. 761–762.
- [4] Masi, M., 2005. "A step beyond tsallis and rényi entropies". *Physics Letters A*, **338**(3-5), pp. 217–224.
- [5] Amabile, T. M., 1996. *Creativity in context: Update to the social psychology of creativity*. Hachette UK.
- [6] Sternberg, R. J., 1999. *Handbook of creativity*. Cambridge University Press.
- [7] Mumford, M. D., and Gustafson, S. B., 1988. "Creativity syndrome: Integration, application, and innovation". *Psychological bulletin*, **103**(1), p. 27.
- [8] Baer, J., 2014. *Creativity and divergent thinking: A task-specific approach*. Psychology Press.
- [9] Torrance, E. P., 1972. "Predictive validity of the torrance tests of creative thinking". *The Journal of creative behavior*, **6**(4), pp. 236–262.
- [10] Acar, S., and Runco, M. A., 2017. "Latency predicts category switch in divergent thinking". *Psychology of Aesthetics, Creativity, and the Arts*, **11**(1), p. 43.
- [11] Dylla, N., 1991. "Thinking methods and procedures in mechanical design". *Mechanical design, technical university of Munich, PhD*.
- [12] Beitz, W., and Pahl, G., 1996. "Engineering design: a systematic approach". *MRS BULLETIN*, **71**.
- [13] Pahl, G., and Beitz, W., 2013. *Engineering design: a systematic approach*. Springer Science & Business Media.
- [14] Dorst, K., and Cross, N., 2001. "Creativity in the design process: co-evolution of problem–solution". *Design studies*, **22**(5), pp. 425–437.
- [15] Henderson, D., Helm, K., Jablokow, K., McKilligan, S., Daly, S., and Silk, E., 2017. "A comparison of variety metrics in engineering design". In ASME 2017 International Design Engineering Technical Conferences and Computers and Information in Engineering Conference, American Society of Mechanical Engineers, pp. V007T06A004–V007T06A004.
- [16] Linsey, J. S., Clauss, E., Kurtoglu, T., Murphy, J., Wood, K., and Markman, A., 2011. "An experimental study of group idea generation techniques: understanding the roles of idea representation and viewing methods". *Journal of Mechanical Design*, **133**(3), p. 031008.
- [17] Oman, S. K., Tumer, I. Y., Wood, K., and Seepersad, C., 2013. "A comparison of creativity and innovation metrics and sample validation through in-class design projects". *Research in Engineering Design*, **24**(1), pp. 65–92.
- [18] Verhaegen, P.-A., Vandevenne, D., Peeters, J., and Duflo, J. R., 2013. "Refinements to the variety metric for idea evaluation". *Design Studies*, **34**(2), pp. 243–263.
- [19] Linsey, J. S., 2007. "Design-by-analogy and representation in innovative engineering concept generation". PhD thesis, University of Texas, Austin.
- [20] Sluis-Thiescheffer, W., Bekker, T., Eggen, B., Vermeeren, A., and De Ridder, H., 2016. "Measuring and comparing novelty for design solutions generated by young children through different design methods". *Design Studies*, **43**, pp. 48–73.
- [21] Peeters, J., Verhaegen, P.-A., Vandevenne, D., and Duflo, J., 2010. "Refined metrics for measuring novelty in ideation". *IDMME Virtual Concept Research in Interaction Design*, Oct, pp. 20–22.
- [22] Fuge, M., Stroud, J., and Agogino, A., 2013. "Automatically inferring metrics for design creativity". *ASME Paper No. DETC2013-12620*.
- [23] Shatz, D., 2004. *Peer review: A critical inquiry*. Rowman & Littlefield.
- [24] Chulvi, V., Mulet, E., Chakrabarti, A., López-Mesa, B., and González-Cruz, C., 2012. "Comparison of the degree of creativity in the design outcomes using different design methods". *Journal of Engineering Design*, **23**(4), pp. 241–269.
- [25] Kline, P., 2014. *The new psychometrics: science, psychology and measurement*. Routledge.
- [26] Twomey, M., Wallis, L. A., and Myers, J. E., 2007. "Limitations in validating emergency department triage scales". *Emergency Medicine Journal*, **24**(7), pp. 477–479.
- [27] Hennessey, B. A., and Amabile, T. M., 1999. "Consensual assessment". In *Encyclopedia of Creativity*, M. Runco and S. Pritzker, eds. Elsevier Science, pp. 347–359.
- [28] Harnad, S., 2008. "Validating research performance metrics against peer rankings". *Ethics in science and environmental politics*, **8**(1), pp. 103–107.
- [29] Baer, J., and McKool, S. S., 2009. "Assessing creativity using the consensual assessment technique". In *Handbook of research on assessment technologies, methods, and applications in higher education*. IGI Global, pp. 65–77.
- [30] Amabile, T. M., and Pillemer, J., 2012. "Perspectives on the social psychology of creativity". *The Journal of Creative Behavior*, **46**(1), pp. 3–15.
- [31] Schaefer, C. E., and Anastasi, A., 1968. "A biographical inventory for identifying creativity in adolescent boys". *Journal of Applied Psychology*, **52**(1p1), p. 42.
- [32] Taylor, C., and Ellison, R., 1966. "Alpha biographical inventory". *Salt Lake City, UT: Institute for Behavioral Research in Creativity*.
- [33] Cropley, A. J., 2000. "Defining and measuring creativity: Are creativity tests worth using?". *Roeper review*, **23**(2), pp. 72–79.
- [34] Douglas, L., Jillian, M., Thomas, L., and Eric, L., 2006. "Identifying quality, novel, and creative ideas: constructs and scales for idea evaluation1". *Journal of the Association for Information Systems*, **7**(10), p. 646.
- [35] Cross, N., and Roy, R., 1989. *Engineering design methods*, Vol. 4. Wiley New York.

- [36] Dow, S., Fortuna, J., Schwartz, D., Altringer, B., Schwartz, D., and Klemmer, S., 2011. "Prototyping dynamics: sharing multiple designs improves exploration, group rapport, and results". In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, Acm, pp. 2807–2816.
- [37] Jansson, D. G., and Smith, S. M., 1991. "Design fixation". *Design studies*, **12**(1), pp. 3–11.
- [38] Kershaw, T. C., and Ohlsson, S., 2004. "Multiple causes of difficulty in insight: the case of the nine-dot problem.". *Journal of experimental psychology: learning, memory, and cognition*, **30**(1), p. 3.
- [39] Wilson, J. O., Rosen, D., Nelson, B. A., and Yen, J., 2010. "The effects of biological examples in idea generation". *Design Studies*, **31**(2), pp. 169–186.
- [40] Srinivasan, V., and Chakrabarti, A., 2010. "Investigating novelty–outcome relationships in engineering design". *AI EDAM*, **24**(2), pp. 161–178.
- [41] Verhaegen, P.-A., Vandevenne, D., Peeters, J., and Duflou, J. R., 2015. "A variety metric accounting for unbalanced idea space distributions". *Procedia engineering*, **131**, pp. 175–183.
- [42] Viswanathan, V. K., and Linsey, J. S., 2012. "Physical models and design thinking: A study of functionality, novelty and variety of ideas". *Journal of Mechanical Design*, **134**(9).
- [43] Thevenot, H. J., and Simpson, T. W., 2007. "A comprehensive metric for evaluating component commonality in a product family". *Journal of Engineering Design*, **18**(6), pp. 577–598.
- [44] Kota, S., Sethuraman, K., and Miller, R., 2000. "A metric for evaluating design commonality in product families". *J. Mech. Des.*, **122**(4), pp. 403–410.
- [45] Jung, S., and Simpson, T. W., 2016. "An integrated approach to product family redesign using commonality and variety metrics". *Research in Engineering Design*, **27**(4), Oct, pp. 391–412.
- [46] Chan, J., Dang, S., and Dow, S. P., 2016. "Comparing different sensemaking approaches for large-scale ideation". In *Proceedings of the 2016 CHI Conference on Human Factors in Computing Systems*, ACM, pp. 2717–2728.
- [47] Chan, J., Fu, K., Schunn, C., Cagan, J., Wood, K., and Kotovsky, K., 2011. "On the benefits and pitfalls of analogies for innovative design: Ideation performance based on analogical distance, commonness, and modality of examples". *Journal of mechanical design*, **133**(8), p. 081004.
- [48] Gini, C., 1912. "Variabilità e mutabilità". *Reprinted in Memorie di metodologica statistica (Ed. Pizetti E, Salvemini, T). Rome: Libreria Eredi Virgilio Veschi.*
- [49] Simpson, E. H., 1949. "Measurement of diversity". *Nature*, **163**(4148), p. 688.
- [50] Shannon, C. E., 1948. "A mathematical theory of communication". *Bell system technical journal*, **27**(3), pp. 379–423.
- [51] Jost, L., 2006. "Entropy and diversity". *Oikos*, **113**(2), pp. 363–375.
- [52] Crupi, V., 2019. "Measures of biological diversity: Overview and unified framework". In *From Assessing to Conserving Biodiversity*. Springer, pp. 123–136.
- [53] Rhoades, S. A., 1993. "The herfindahl-hirschman index". *Fed. Res. Bull.*, **79**, p. 188.
- [54] Greenberg, J. H., 1956. "The measurement of linguistic diversity". *Language*, **32**(1), pp. 109–115.
- [55] Pitt, R., Pirtle, W. N. L., and Metzger, A. N., 2019. "Academic specialization, double majoring, and the threat to breadth in academic knowledge". *The Journal of General Education*, **66**(3-4), pp. 166–191.
- [56] Nelson, J. D., Crupi, V., Meder, B., Cevolani, G., and Tentori, K., 2017. "A unified model of entropy and the value of information". *Decision Making*, **7**, pp. 119–148.
- [57] Ahmed, F., Ramachandran, S. K., Fuge, M., Hunter, S., and Miller, S., 2019. "Measuring and optimizing design variety using herfindahl index". In *ASME 2019 International Design Engineering Technical Conferences and Computers and Information in Engineering Conference*, American Society of Mechanical Engineers Digital Collection.
- [58] Hurlbert, S. H., 1971. "The nonconcept of species diversity: a critique and alternative parameters". *Ecology*, **52**(4), pp. 577–586.
- [59] Nemhauser, G. L., Wolsey, L. A., and Fisher, M. L., 1978. "An analysis of approximations for maximizing submodular set functions". *Mathematical Programming*, **14**(1), pp. 265–294.
- [60] Feige, U., Mirrokni, V. S., and Vondrak, J., 2011. "Maximizing non-monotone submodular functions". *SIAM Journal on Computing*, **40**(4), pp. 1133–1153.
- [61] Krause, A., and Golovin, D., 2014. "Submodular function maximization". In *Tractability: Practical Approaches to Hard Problems*. Cambridge University Press, pp. 71–104.
- [62] Hoffmann, S., et al., 2008. Generalized distribution based diversity measurement: Survey and unification. Tech. rep., Otto-von-Guericke University Magdeburg, Faculty of Economics and Management.
- [63] Tsallis, C., 1988. "Possible generalization of boltzmann-gibbs statistics". *Journal of statistical physics*, **52**(1-2), pp. 479–487.
- [64] Keylock, C., 2005. "Simpson diversity and the shannon–wiener index as special cases of a generalized entropy". *Oikos*, **109**(1), pp. 203–207.
- [65] Stobbe, P., and Krause, A., 2010. "Efficient minimization of decomposable submodular functions". In *Advances in Neural Information Processing Systems*, pp. 2208–2216.
- [66] Kendall, M., 1962. *Rank correlation methods*. Theory and applications of rank order-statistics. Hafner Pub. Co.
- [67] Stewart, N., Brown, G. D., and Chater, N., 2005. "Absolute identification by relative judgment.". *Psychological review*, **112**(4), p. 881.
- [68] Ahmed, F., Fuge, M., and Gorbunov, L. D., 2016. "Discovering diverse, high quality design ideas from a large

- corpus”. In ASME International Design Engineering Technical Conferences, ASME.
- [69] Ahmed, F., and Fuge, M., 2018. “Ranking ideas for diversity and quality”. *Journal of Mechanical Design*, **140**(1), p. 011101.
 - [70] Starkey, E. M., Hunter, S. T., and Miller, S. R., 2019. “Are creativity and self-efficacy at odds? an exploration in variations of product dissection in engineering education”. *Journal of Mechanical Design*, **141**(1), p. 012001.
 - [71] Toh, C. A., and Miller, S. R., 2014. “The impact of example modality and physical interactions on design creativity”. *Journal of Mechanical Design*, **136**(9).
 - [72] Ahmed, F., Ramachandran, S. K., Fuge, M., Hunter, S., and Miller, S., 2019. “Interpreting idea maps: Pairwise comparisons reveal what makes ideas novel”. *Journal of Mechanical Design*, **141**(2), p. 021102.
 - [73] Ritchey, T., 2006. “Problem structuring using computer-aided morphological analysis”. *Journal of the Operational Research Society*, **57**(7), pp. 792–801.
 - [74] Pahl, G., Beitz, W., Feldhusen, J., and Grote, K., 2007. *Engineering Design: A Systematic Approach*. Solid mechanics and its applications. Springer London.
 - [75] George, D., Linnerud, B., and Mocko, G., 2014. “Integrated idea generation method for concept generation using morphological and options matrices”. In Proc. TMCE, pp. 1–12.