



The
tkl-Score
tkl-Score for Data-Sharing Misuseability
Kalvin Eng, Eleni Stroulia

► To cite this version:

Kalvin Eng, Eleni Stroulia. The

tkl-Score

tkl-Score for Data-Sharing Misuseability. 35th IFIP Annual Conference on Data and Applications Security and Privacy (DBSec), Jul 2021, Calgary, AB, Canada. pp.312-324, 10.1007/978-3-030-81242-3_18. hal-03677040

HAL Id: hal-03677040
<https://inria.hal.science/hal-03677040>

Submitted on 24 May 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution 4.0 International License



This document is the original author manuscript of a paper submitted to an IFIP conference proceedings or other IFIP publication by Springer Nature. As such, there may be some differences in the official published version of the paper. Such differences, if any, are usually due to reformatting during preparation for publication or minor corrections made by the author(s) during final proofreading of the publication manuscript.

The *tkl*-Score for Data-Sharing Misuseability

Kalvin Eng and Eleni Stroulia

University of Alberta, Edmonton AB T6G 2R3, Canada
{kalvin.eng, stroulia}@ualberta.ca

Abstract. Estimating the potential for data misuse is essential for all data-sharing decisions. This work presents *tkl*-Score which extends the state of the art *M*-Score and *L*-Severity measures. The new proposed measure is sensitive to the increased misuse potential when records are more identifiable in a source table with *l*-Distinguishing Factor and also when sensitive attributes are less granular in a source table with *l*-Distinguishing Factor and *t*-Distinguishing Factor; in contrast, the earlier *M*-Score and *L*-Severity only account for record identifiability in a source table with *k*-Distinguishing Factor. *tkl*-Score is shown to better characterize the risk of releasing records compared to *M*-Score and *L*-Severity due to accounting for sensitive attribute granularity.

Keywords: Data Misuse · Data Sharing · Data Privacy.

1 Introduction

As data-collection software becomes more ubiquitous, due to mobile apps, social platforms, and the internet of things (IoT), the privacy and sensitivity of collected data is becoming an increasing concern to data stakeholders. This ever-increasing amount of devices and apps leads to a plethora of data being collected. Organizations who wish to share their data need to be more aware of the consequences that can affect the different stakeholders of the data being shared. Unnecessary disclosures of sensitive information can lead to severe consequences for the subjects of the data and legal repercussions to the organizations themselves. For example, the personal information of over 50 million Facebook users were unintentionally exposed to Cambridge Analytica and used for political gain leading to backlash against Facebook [1].

The objective of our work is to develop a misusability score that when given a dataset to be shared, the score accurately estimates the risk of potential misuse of this data and helps inform the decision making of the owner organization. The rest of this paper is organized as follows. In Section 2, we review the current state of misuseability metrics and highlight their drawbacks that we wish to address. Based on these drawbacks, we introduce *tkl*-Score and its derivative *tkl*-Score_{max} in Section 3 which augments current misuseability metrics. In Section 4, we perform a comparative analysis between *tkl*-Score and previous misuseability scores highlighting why previous scores should be augmented. Finally in Section 5, we summarize our work and introduce future avenues of exploration.

2 Related Work on Misuseability Scores

We base our work on *M-Score* and *L-Severity* which are two key metrics for estimating the potential of data misuse. Both are calculated using weights derived from a “source table” to help estimate the risk of potential misuseability in a subset of a source table called a “published table”. The scores also rely on two types of attributes. *Quasi-identifiers* which are attributes that, when combined with other information, may partially reveal an individual’s identity. For example, when two datasets have “city of residence” and “gender” columns, their join based on these quasi-identifier attributes may reveal identifying information such as an individual’s name present only in one of the datasets. As well, there are *sensitive attributes* which convey information that should not be exposed publicly (e.g. one’s health condition), because the release of this information has the potential to harm an individual such as damaging their reputation.

2.1 *M-Score*

M-Score [2] is the first metric designed specifically for identifying the potential negative impacts of a dataset release. It is a score for tabular data that quantifies the ability of a user to maliciously exploit exposed data and takes into account the anonymity of individuals in a dataset as well as the sensitivity of the data attribute values. The process of calculating the *M-Score* for a published table consists of two steps: (1) eliciting weights for sensitive attribute values, and (2) calculating the *M-Score* for the published table.

(1) Eliciting Weights of Sensitive Attribute Values There are several methods for eliciting sensitive value weights from domain experts, but the authors of *M-Score* argue that the Analytic Hierarchy Process [5] elicits the best results for discretized data.

(2) *M-Score* of the Published Table Given the weights of the sensitive attribute values, a source table, and a published table (i.e. a data subset of the source table), *M-Score* can be calculated. To begin, each record of a published table is given a record score as follows:

$$RS_{M_r} = \frac{\min(1, \sum_{A_{S_i} \in r} \text{weight}(A_{S_i}))}{DF_{k_r}} \quad (1)$$

A record score RS_{M_r} of the r th record of a published table is the sum of each i th sensitive attribute value weight of the record minimized to 1 divided by the k -Distinguishing Factor DF_{k_r} of the r th record.

The k -Distinguishing Factor is a measure dependent on comparing the published table to the source table. It quantifies how easily an individual can be identified, based on the distinctiveness (or uniqueness) of records in a “lookup table”. The “lookup table” is collection of records related to a population that can identify an individual. Since such a collection is not easily acquired, the “lookup table” is approximated to be the source table. It is based on the k -anonymity measure [6] that groups rows with similar quasi-identifier values into “equivalence classes”, where an equivalence class is a set of records that have the same values for quasi-identifiers attributes [3].

The k -Distinguishing Factor of a record in a published table is the size of the equivalence class in the source table that contains the record in the published table. If there are no quasi-identifiers to form an equivalence class, then the k -Distinguishing Factor of a record is the size of the published table. The k -Distinguishing Factor is meant to account for how distinguishable an exposed record is when it is published from the source table, and helps to differentiate the records by their identifiability — when the k -Distinguishing Factor is smaller, the record is more identifiable and the record score therefore becomes larger.

With all record scores of a published table calculated, the M -Score of a published table can be computed based on the maximal record score and the number of records, n , in the published table.

$$M\text{-Score} = n^{\frac{1}{x}} \times \max_{0 \leq r \leq n} \left(\frac{\min(1, \sum_{A_{S_i} \in} \text{weight}(A_{S_i}))}{DF_{k_r}} \right) \quad (2)$$

where $x \geq 1$ is a parameter for the importance of the amount of records, A_{S_i} is the i th sensitive attribute value of a record r , and DF_{k_r} is the k -Distinguishing Factor of a record r .

Using the definition above, the M -Score of the published table can be computed by multiplying the highest individual record score among the n records weighted with a power $\frac{1}{x}$. The x of $n^{\frac{1}{x}}$ is a parameter for specifying the importance of the quantity of records in a published table. If $x = 1$ then the amount of records is given more importance compared to the sensitivity of data. If $x \rightarrow \infty$, then $n^{\frac{1}{\infty}} \approx 1$ which means that we would like to know the highest individual record score of M -Score. The parameter x can be assigned any value where $x \geq 1$ with a trade-off between the importance of the highest individual record score to the importance of the amount of records.

Drawbacks M -Score is approximative in nature as it takes the maximum record score of a published table for its score calculation. For example, consider a source table where 99 records of 100 records have a record score of 0.0001 with the remaining record having a record score of 1. A published table with nine records with score 0.0001 and one record with score 1 will result in the same M -Score, i.e., 10, as a published table with nine records of 1 and one of 0.0001.

The approximation in M -Score leads to issues when attempting to identify the “percentage of severity” a published table takes from the source table score. If we consider the example from the previous paragraph, where the published table has a score of 10, and the source table has a score of 100, we can say that the published table makes up 10 percent of the severity of the source table: $10/100 = 0.1$. However, this is not the case when we compare the sum of the actual values of the record scores in each of the tables: $1.0009/1.0099 = 0.991$.

It can also be difficult to decide on the parameter x to model the trade-off between the importance of the number of records or the maximum record score in a published table. If we would like to approximate the severity based on the amount of records, then x can be set as 1. If we assume that releasing any single

maximum record of a published table is the maximum severity, then x can be set as $x \rightarrow \infty$. However, to decide on a value between 1 and $x \rightarrow \infty$, that represents the trade-off between these two factors, an ad hoc decision would need to be made.

The last drawback to note is that the k -Distinguishing Factor in M -Score only accounts for identity disclosure attacks. The authors of M -Score suggest that measures such as l -diversity can be used to account for attribute disclosure attacks, but provide no method to do so.

2.2 L -Severity

Building on the work of M -Score, Vavilis et al. [7] designed L -Severity, a misuse-ability score aiming to address the approximative nature of the M -Score calculation. The process of calculating L -Severity can be divided into three steps: (1) developing a data model of the sensitive attributes in a source table, (2) eliciting weights for sensitive attribute values from the data model, and (3) calculating the overall L -Severity for the published table.

(1) Developing the Data Model The data model in L -Severity is designed to represent the hierarchy of concepts surrounding the sensitive attributes in a source table. Each sensitive attribute of a source table is represented as a node and can fall under more general nodes assigned by a domain expert. As well, the concept of “inference relationships” which link sensitive attributes is also introduced to highlight how sensitive attributes may be related.

(2) Eliciting Sensitive Attribute Value Weights using the Data Model Using the data model, domain experts assign “sensitivity values” to highlight the importance of nodes in the data model as well as “inference values” to quantify the importance between any relationships of attributes in the data model. Once all values have been assigned, they can be used to calculate the sensitive attribute value weights needed for L -Severity.

(3) L -Severity of the Published Table To determine the L -Severity of a published table, the record scores are summed together. A record score can be calculated as follows:

$$RS_{L_r} = \frac{\sum_{A_{S_i} \in r} \text{weight}(A_{S_i})}{DF_{k_r}} \quad (3)$$

A record score RS_{L_r} of the r th record of a published table is the sum of each i th sensitive attribute value weight of the record divided by the k -Distinguishing Factor DF_{k_r} of the r th record. Then, given a published table T , the L -Severity of the table is:

$$L\text{-Severity} = \sum_{r \in T} \left(\frac{\sum_{A_{S_i} \in r} \text{weight}(A_{S_i})}{DF_{k_r}} \right) \quad (4)$$

where r is each record of the published table, A_{S_i} is the i th sensitive attribute value of a record r , and DF_{k_r} is the k -Distinguishing Factor of a record r .

Drawbacks Although L -Severity addresses the approximative nature of M -Score by summing record scores, L -Severity cannot account for the case that assumes releasing the maximum record score of a published table is the maximum severity. Since L -Severity does not have a x parameter like M -Score, it cannot control the trade-off between the importance of the number of records and the maximum severity to calculate its score. Instead, it only accounts for the case that calculates severity based on the number of published records.

L -Severity also suffers the same drawback as M -Score, failing to consider anonymity measures like l -diversity and t -closeness to account for attribute disclosure. Only identity disclosure is accounted for with k -Distinguishing Factor. The authors of L -Severity suggest that attribute disclosure measures can be integrated into misuseability scoring, but provide no method to do so.

3 Calculating tkl -Score

Expanding on these earlier measures, we propose tkl -Score which incorporates the privacy preserving data publishing metrics: l -diversity as l -Distinguishing Factor, and t -closeness as t -Distinguishing Factor to account for attribute disclosure attacks in addition to identity disclosure attacks with k -Distinguishing Factor. It is important to also account for attribute disclosure as it may be difficult to be certain of an individual's identity, but attributes relating to an individual can still be disclosed when similar information about identities are grouped together.

To demonstrate how l -Distinguishing Factor and t -Distinguishing Factor are used in tkl -Score, we introduce the record score of tkl -Score which is defined as:

$$RS_{tkl_r} = \frac{DF_{t_r} + \sum_{A_{S_i} \in r} \text{weight}(A_{S_i})}{DF_{l_r}} \quad (5)$$

where r is a record, DF_{l_r} is the l -Distinguishing Factor of a record, DF_{t_r} is the t -Distinguishing Factor of a record, and $\text{weight}(A_{S_i})$ is the weight of the i th sensitive attribute value of a record.

It should be noted that l -Distinguishing Factor and t -Distinguishing Factor, like k -Distinguishing Factor, aims to quantify the identity and attribute uniqueness of records in a “lookup table” that contains all records related to a population. However, since it is difficult to obtain all records related to a population, the “lookup table” is approximated to be the source table of a published table.

3.1 l -Distinguishing Factor

l -diversity [4] is the measure used to determine how distinguishable an individual is based on attribute frequency in an equivalence class (a group of records with common quasi-identifier attribute values). In this paper, we use the l -diversity definition where every equivalence class in a table contains at least l distinct values for a sensitive attribute in order to be l -diverse. An interesting property of l -diversity is that l will always be $\leq k$ of k -anonymity.

Proof. Let k be the k -anonymity of a dataset which is the smallest equivalence class with size k rows. Assume that the l -diversity of the smallest equivalence

class has $k < l$. Based on the assumption, let $l = k + 1$. Then based on the definition of l -diversity, there must be $k + 1$ unique attribute values. However, this is a contradiction as the size of the equivalence class must be $k + 1$ to have $k + 1$ unique attribute values. Now assume that the dataset has multiple sensitive attributes and therefore multi-attribute l -diversity [4] is used to create different grouping combinations. The largest equivalence class of the combinations will still be at most be the size of the original equivalence class that is matched with only quasi-identifiers. I.e. as more attributes need to be matched to form a grouping of records, the size of the groupings either remains the same as the attribute values are all the same, or the size of the grouping becomes smaller when the attribute values are different. Therefore, $l \leq k$.

Recall that k -Distinguishing Factor is used to determine how distinguishable a record is in a source table and is a divisor in the score equations of M -Score and L -Severity. A large k -Distinguishing Factor implies a lower risk of identifiability, as more records are required to uniquely identify an individual. Hence, we wish to capture the maximal severity by minimizing k -Distinguishing Factor.

Since $l \leq k$, the k -anonymity metric for identity disclosure attacks of k -Distinguishing Factor will always be accounted for when using l -Distinguishing Factor. Therefore in tkl -Score, l -Distinguishing Factor replaces the k -Distinguishing Factor factor used in M -Score and L -Severity.

The definition of l -Distinguishing Factor uses the following definition of multi-attribute l -diversity: Let T be a table with nonsensitive attributes $Q_1, \dots, Q_{m_1}$ and sensitive attributes $S_1, \dots, S_{m_2}$. If for all iterations $i = 1 \dots m_2$, the table T is l -diverse when S_i is treated as the sole sensitive attribute and $\{Q_1, \dots, Q_{m_1}, S_1, \dots, S_{i-1}, S_{i+1}, \dots, S_{m_2}\}$ is treated as the “quasi-identifiers” to form “equivalence classes” [4].

The l -Distinguishing Factor of a record in a published table is the minimal multi-attribute l -diversity [4] “equivalence class” in the source table that contains the record in the published table. If there are no quasi-identifiers to form an equivalence class in the source table, then the minimal multi-attribute l -diversity of the published table (forming equivalence classes based on the sensitive attributes only) is the l -Distinguishing Factor of a record.

3.2 t -Distinguishing Factor

t -closeness [3] provides a measurement for the similarity between the attribute value distribution of an equivalence class and the attribute value distribution of an entire table. The similarity of distributions for sensitive attribute values helps to determine the true diversity of sensitive attributes globally. In comparison, l -diversity only considers the diversity of sensitive attributes within an equivalence class, which means sensitive attributes may still be prone to attribute disclosure if all the unique attributes are in a single equivalence class. As a result, t -closeness and l -diversity are two different anonymity measures that should be used together as a way to measure the anonymity of sensitive attributes in a table.

The objective of integrating t -closeness into tkl -Score is to increase the relative severity of an exposure as the value of the t -closeness of the released records increases. This is because the higher the t -closeness, the higher the likelihood of an attribute disclosure attack as the distributions are less similar and sensitive attributes are easier to discern. Therefore, t is added to the record score.

We assume to have a record score function $\frac{S}{DF}$ modeled after Equation 3 where $1 \geq S \geq 0$ that represents the sum of sensitive attribute value weights of a record, and $DF \geq 1$ that represents the l -Distinguishing Factor. The addition of t to the numerator of the function to calculate a record score produces a linear translation of the score by $\frac{t}{DF}$ on the y-axis when visualized on a plot. In the best case the record score is minimally reduced when $t \approx 0$ (less risk because the attribute distributions are similar throughout the table) and in the worst case when $t = 1$ the score increases by $\frac{t}{DF}$ (more risk since equivalence class attribute distributions are not similar to the global attribute distribution). As a result, the record score will either be maintained or increased by a proportional t factor to indicate the severity of a record release more finely.

The t -Distinguishing Factor of a record in a published table is the maximal t -closeness of the equivalence class in the source table that contains this record. If there are no quasi-identifiers to form an equivalence class in the source table, then the t -Distinguishing Factor of a record is 0 as there are no equivalence classes to compare the distribution of sensitive attribute values. This is consistent with the limitation of k -anonymity which would also be 0 when there are no quasi-identifiers to form an equivalence class and indicates that there are no attributes that can be referenced to identify a person.

3.3 tkl -Score

To calculate the tkl -Score of a published table, every record in the table is scored and summed together using sensitive attribute value weights. To obtain sensitive attribute value weights for tkl -Score, the same methods used to derive weights in M -Score and L -Severity can be used. Equation (6) defined below uses the record score defined in Equation (5).

Given a published table T , the tkl -Score is calculated as:

$$tkl\text{-Score} = \sum_{r \in T} \left(\frac{DF_{t_r} + \sum_{A_{S_i} \in r} \text{weight}(A_{S_i})}{DF_{l_r}} \right) \quad (6)$$

where for each record r , the total sum of each i th sensitive attribute value weight A_{S_i} is summed with the t -Distinguishing Factor DF_{t_r} of r and divided with the l -Distinguishing Factor DF_{l_r} of r .

We also introduce $tkl\text{-Score}_{\max}$ which is a score that is modeled after M -Score ($x \rightarrow \infty$) to signify that releasing any one maximum record score is of maximum severity. The difference of $tkl\text{-Score}_{\max}$ and M -Score ($x \rightarrow \infty$) is that it also accounts for attribute disclosure because it incorporates l -Distinguishing Factor and t -Distinguishing Factor, instead of only k -Distinguishing Factor.

Given a table with n records, the table's $tkl\text{-Score}_{\max}$ is then:

$$tkl\text{-Score}_{\max} = \max_{0 \leq r \leq n} \left(\frac{DF_{t_r} + \sum_{A_{S_i} \in r} \text{weight}(A_{S_i})}{DF_{l_r}} \right) \quad (7)$$

where for each record r , the total sum of each i th sensitive attribute value $\text{weight}(A_{S_i})$ is summed with the t -Distinguishing Factor DF_{t_r} of r and divided with the l -Distinguishing Factor DF_{l_r} of r .

4 Comparative Analysis

To compare $tkl\text{-Score}$ against its predecessors $M\text{-Score}$ and $L\text{-Severity}$, three alternative normalization assumptions can be considered: (1) the maximum severity is when a complete source table is released, (2) the maximum severity is when any one record with the maximum record score in the source table is released, and (3) the maximum severity is when a score reaches the theoretical maximum score determined by bounding the sum of sensitive attribute value weights and distinguishing factors. With assumption (1) the severity is related to the number of released records and therefore $tkl\text{-Score}$, $L\text{-Severity}$, and $M\text{-Score}$ ($x = 1$) are compared. With assumption (2) the severity is related to the maximum record of a source table, and therefore $tkl\text{-Score}_{\max}$ and $M\text{-Score}$ ($x \rightarrow \infty$) are compared. For assumption (3) any misuseability score can be used, but it is best used with $tkl\text{-Score}_{\max}$ and $M\text{-Score}$ ($x \rightarrow \infty$) to compare their severity estimates.

Harel et al. [2] suggest that $M\text{-Score}$ can be normalized under assumption (1) by taking the $M\text{-Score}$ of the published table and dividing by the $M\text{-Score}$ of the source table. This form of normalization assumes that releasing a complete source table is the maximum severity — when a subset of a source table is published, it will take a percentage of the source table score.

In this section, we illustrate how t -Distinguishing Factor and l -Distinguishing Factor in $tkl\text{-Score}$ affects the row scores of two different tables using assumption (1) for normalizing the misuseability scores to the $[0, 1]$ range for comparison. Figure 1 illustrates how the misuseability scores can score rows differently based on attributes in a table.

4.1 A Case for t -Distinguishing Factor

The t -Distinguishing Factor is used to determine a record's distinctiveness based on the similarity of the distribution of sensitive attributes in the equivalence class of the source table containing this record and the distribution of sensitive attributes in the whole table. Recall that the higher the t value, the higher the severity a record would be as the distribution of the values of the sensitive attributes and the whole table are more distinct and therefore easier to differentiate.

To illustrate how t -Distinguishing Factor helps to distinguish records, the record scores for the source Table 1a are calculated in Table 1d. From the normalized scores in Table 1b, we can see that: (i) $tkl\text{-Score}$ is reduced from row 5

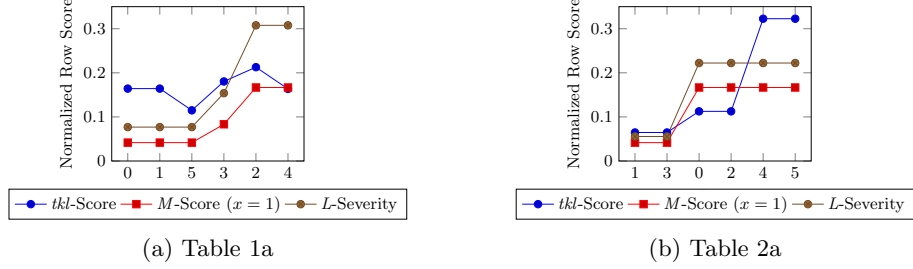


Fig. 1: Per row scores normalized against the source table score for tkl -Score, M -Score ($x = 1$), and L -Severity of the records in Table 1a and Table 1a. It is sorted ascending by L -Severity.

to row 1, while M -Score ($x = 1$) and L -Severity maintains consistency for rows 0, 1, and 5; and (ii) the tkl -Score is reduced from row 2 to row 4, while M -Score ($x = 1$) and L -Severity are increased. These observations can be visualized in Figure 1a.

Note that the M -Score (regardless of x parameter) and L -Severity calculated for each record seen in Table 1d are the same because there is only a single record for the misuseability score calculation. However, if we were to calculate the misuseability scores for a larger subset of published records, M -Score ($x = 1$), M -Score ($x \rightarrow \infty$) and L -Severity will produce different results.

The difference between (i) and (ii) stems from the equivalence class of rows 4 and 5 in Table 1a. The distribution of the *Initial Diagnosis* values in this equivalence has a lower t value compared to the other equivalence classes of the table as “HIV” and “Migraine” occur elsewhere in the table. In contrast, the equivalence class with rows 0 and 1, and the equivalence class with rows 2 and 3 have a unique value that does not appear in any other equivalence class leading to a higher t value. Therefore, rows 4 and 5 have a lower t -Distinguishing Factor than the other records, and as a result, rows 4 and 5 are considered less severe by tkl -Score than the other records in the table that have the same l -Distinguishing Factor and sum of sensitive attribute value weights as seen in Table 1d.

Intuitively, observing the attribute values of Table 1a, we can see that either a female lawyer from Edmonton or a male lawyer from Edmonton has “HIV”. Likewise, either a female lawyer from Edmonton or a female lawyer from Calgary has a “Migraine”. Since “HIV” and “Migraine” must be distinguished from two distinct equivalence classes, it is less likely that these sensitive attributes will be inferred as opposed to “Flu” and “Hypertension” which occur only in a single equivalence class.

4.2 A Case for l -Distinguishing Factor

The l -Distinguishing Factor is used to determine a record’s distinctiveness based on the size of the equivalence class that contains this record in the source table and also the values of sensitive attributes in the equivalence class that contains

id	Job	City	Gender	Initial Diagnosis
0	Lawyer	Calgary	Female	Flu
1	Lawyer	Calgary	Female	Migraine
2	Lawyer	Edmonton	Male	HIV
3	Lawyer	Edmonton	Male	Hypertension
4	Lawyer	Edmonton	Female	HIV
5	Lawyer	Edmonton	Female	Migraine

(a) The sensitive attribute of this source table is *Initial Diagnosis*; the quasi-identifiers are *Job*, *City*, and *Gender*.

Row	<i>tkl</i> -Score	<i>M</i> -Score ($x = 1$)	<i>L</i> -Severity
0	0.1642	0.0417	0.0769
1	0.1642	0.0417	0.0769
5	0.1149	0.0417	0.0769
3	0.1804	0.0833	0.1538
2	0.2128	0.1667	0.3077
4	0.1635	0.1667	0.3077

(b) Misuseability scores for each row of Table 1a normalized against the source table score rounded to four decimal places.

Disease	Medication	Age	Initial Diagnosis
$W(Flu) = 0.0864$	$W(Antibiotics) = 0.10432$	$W(30+) = 0.08$	$W(Migraine) = 0.05472$
$W(H1N1) = 0.3456$	$W(Paracetamol) = 0.10432$	$W(< 30) = 0.08$	$W(Flu) = 0.05472$
$W(Hypertension) = 0.3456$	$W(ARV) = 0.10432$		$W(Hypertension) = 0.10944$
$W(HIV) = 0.432$	$W(Tamiflu) = 0.10432$		$W(HIV) = 0.21888$
	$W(Statin) = 0.10432$		

(c) Sensitive Attribute Value Weights

Row	<i>tkl</i> -Score	<i>M</i> -Score	<i>L</i> -Severity	DF_t	DF_k	DF_l	Weights
0	0.27736	0.02736	0.02736	0.50000	2	2	0.05472
1	0.27736	0.02736	0.02736	0.50000	2	2	0.05472
5	0.19403	0.02736	0.02736	0.33333	2	2	0.05472
3	0.30472	0.05472	0.05472	0.50000	2	2	0.10944
2	0.35944	0.10944	0.10944	0.50000	2	2	0.21888
4	0.27611	0.10944	0.10944	0.33333	2	2	0.21888

(d) For each row of Table 1a: the raw misuseability scores, distinguishing factors, and sum of sensitive attribute value weights (Table 1c) rounded to five decimal places.

Table 1: A case for using *t*-Distinguishing Factor.

this record in the source table. Recall that *l*-Distinguishing Factor is part of the denominator of a record score (Equation 5), and therefore the smaller the *l*-Distinguishing Factor the more the potential implications of releasing a record. So if there are more unique sensitive attribute values in an equivalence class, the harder it will be to link specific sensitive attributes to the identities of an equivalence class.

To illustrate how *l*-Distinguishing Factor can distinguish records, the record scores for the source Table 2a are calculated in Table 2c. From the normalized scores in Table 2b, we can see that *M*-Score and *L*-Severity maintain consistent scores for rows 0, 2, 4, and 5 while *tkl*-Score has greater scores in rows 4-and-5 compared to rows 0-and-2. This observation can be visualized in Figure 1b.

We can account for the discrepancies between rows 4-and-5 and rows 0-and-2 in *tkl*-Score by observing the corresponding *l*-Distinguishing Factor and *t*-Distinguishing Factor values. In Table 2c, we see that the *t*-Distinguishing Factor values of rows 4-and-5 are twice as much as the *t*-Distinguishing Factor values of rows 0-and-2. This increases *tkl*-Score additively, but does not explain why the *tkl*-Score of rows 4-and-5 are more than double rows 0-and-2. To explain this multiplicative effect, we examine the *l*-Distinguishing Factor values of the rows.

id	Job	City	Gender	Initial Diagnosis
0	Lawyer	Calgary	Female	HIV
1	Lawyer	Calgary	Female	Flu
2	Lawyer	Edmonton	Male	HIV
3	Lawyer	Edmonton	Male	Flu
4	Lawyer	Edmonton	Female	HIV
5	Lawyer	Edmonton	Female	HIV

(a) The sensitive attribute of this source table is *Initial Diagnosis*; the quasi-identifiers are *Job*, *City*, and *Gender*.

Row	tkl -Score	M -Score ($x = 1$)	L -Severity
1	0.0647	0.0417	0.0556
3	0.0647	0.0417	0.0556
0	0.1126	0.1667	0.2222
2	0.1126	0.1667	0.2222
4	0.3227	0.1667	0.2222
5	0.3227	0.1667	0.2222

(b) Misuseability scores for each row of Table 2a normalized against the source table rounded to four decimal places.

Row	tkl -Score	M -Score	L -Severity	DF_t	DF_k	DF_l	Weights
1	0.11069	0.02736	0.02736	0.16667	2	2	0.05472
3	0.11069	0.02736	0.02736	0.16667	2	2	0.05472
0	0.19277	0.10944	0.10944	0.16667	2	2	0.21888
2	0.19277	0.10944	0.10944	0.16667	2	2	0.21888
4	0.55221	0.10944	0.10944	0.33333	2	1	0.21888
5	0.55221	0.10944	0.10944	0.33333	2	1	0.21888

(c) For each row of Table 2a: the raw misuseability scores, distinguishing factors, and sum of sensitive attribute value weights (Table 1c) rounded to five decimal places.

Table 2: A case for using l -Distinguishing Factor

From the equivalence class containing rows 4-and-5 as seen in Table 2a, it can be observed that “HIV” is the only unique sensitive attribute in this equivalence class. In every other equivalence class, there are two unique values. This means that if we knew a female lawyer was from Edmonton, we can deduce that they have HIV if the table were to be released, as opposed to having to distinguish between two different sensitive attribute values with the other equivalence class groupings. Because the equivalence class containing rows 4-and-5 has only “HIV” as the sensitive attribute value, the l -Distinguishing Factor for rows 4-and-5 is 1. As a result, the record scores of rows 4-and-5 are not reduced by a factor of 2 like rows 0-and-2.

4.3 Limitations

The basis of misuseability scores are the sensitive attribute value weights defined by domain experts. Therefore the “misuseability” of data is entirely dependent on how attributes are classified. This drawback is no different than M -Score or L -Severity, but our work shows how misuseability using pre-existing classifications can be more finely determined when integrating additional anonymity measures.

As well, misuseability scores normalized against one source table should not be compared to a score normalized against a different source table. The reason is that source tables can have different sensitivities of attributes, vary in size, and have different attribute values meaning that the records may have different distinguishing factors. However, misuseability scores are beneficial when deciding the misuseability of a subset of records in a source table that have been/will be shared (i.e. which release of a subset may be more risky than another release).

5 Conclusion

In this paper, we extend previous misuseability scoring by incorporating new distinguishing factors to account for sensitive attribute distinctiveness. We demonstrate how l -Distinguishing Factor and t -Distinguishing Factor in tkl -Score can account for data misuse scenarios not detected by previous misuseability scores. Two cases are presented to demonstrate how t -Distinguishing Factor and l -Distinguishing Factor can distinguish records better than k -Distinguishing Factor and make record scores more granular. Because of these new distinguishing factors, tkl -Score and tkl -Score_{max} are better at characterizing the severity of records compared to L -Severity and M -Score.

Misuseability scoring enables comparisons between different datasets concerning the severity of a release by quantifying the sensitivity of data. Our tkl -Score is presented as an improvement to existing misuseability scoring by accounting for attribute disclosure attacks in addition to identity disclosure attacks.

Future work includes investigating systematic methodologies for determining sensitive attribute weights to address the limitation of misuseability scores being dependent on the classification of attributes. As well, applications of tkl -Score and tkl -Score_{max} should also be investigated for how they might be used to quantify the sensitivity of records leaked in a data breach and indicate the extent of a breach. Furthermore, tkl -Score and tkl -Score_{max} could be integrated with data-loss prevention systems to monitor for any anomalies in user behaviour when accessing database data by calculating a score each time data is accessed by a user and identifying any extreme scores. tkl -Score and tkl -Score_{max} could also be used as part of a risk-assessment process to determine, based on a score, which sensitive records could cause issues when released.

References

1. Graham-Harrison, E., Cadwalladr, C.: Revealed: 50 million Facebook profiles harvested for Cambridge Analytica in major data breach. The Guardian (2018), <https://www.theguardian.com/news/2018/mar/17/cambridge-analytica-facebook-influence-us-election>
2. Harel, A., Shabtai, A., Rokach, L., Elovici, Y.: M-score: A Misuseability Weight Measure. IEEE Transactions on Dependable and Secure Computing **9**(3), 414–428 (2012)
3. Li, N., Li, T., Venkatasubramanian, S.: t-closeness: Privacy beyond k-anonymity and l-diversity. In: 2007 IEEE 23rd International Conference on Data Engineering. pp. 106–115. IEEE (2007)
4. Machanavajjhala, A., Kifer, D., Gehrke, J., Venkatasubramanian, M.: l-diversity: Privacy beyond k-anonymity. ACM Transactions on Knowledge Discovery from Data (TKDD) **1**(1), 3-es (2007)
5. Saaty, T.L.: A Scaling Method for Priorities in Hierarchical Structures. Journal of Mathematical Psychology **15**(3), 234–281 (1977)
6. Sweeney, L.: k-anonymity: A model for protecting privacy. International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems **10**(05), 557–570 (2002)
7. Vavilis, S., Petković, M., Zannone, N.: A Severity-based Quantification of Data Leakages in Database Systems. Journal of Computer Security **24**(3), 321–345 (2016)