

UCSF

UC San Francisco Previously Published Works

Title

Is seizure frequency variance a predictable quantity?

Permalink

<https://escholarship.org/uc/item/5787h035>

Journal

Annals of Clinical and Translational Neurology, 5(2)

ISSN

2328-9503

Authors

Goldenholz, Daniel M

Goldenholz, Shira R

Moss, Robert

et al.

Publication Date

2018-02-01

DOI

10.1002/acn3.519

Copyright Information

This work is made available under the terms of a Creative Commons Attribution-NonCommercial-NoDerivatives License, available at

<https://creativecommons.org/licenses/by-nc-nd/4.0/>

Peer reviewed

RESEARCH ARTICLE

Is seizure frequency variance a predictable quantity?

Daniel M. Goldenholz^{1,2} , Shira R. Goldenholz², Robert Moss³, Jacqueline French⁴, Daniel Lowenstein⁵, Ruben Kuzniecky⁴, Sheryl Haut⁶, Sabrina Cristofaro⁴, Kamil Detyniecki⁷, John Hixson⁵, Philippa Karoly⁸, Mark Cook⁸, Alex Strashny⁹ & William H. Theodore¹

¹Clinical Epilepsy Section, NINDS, NIH, Bethesda, Maryland 20892

²Beth Israel Deaconess Medical Center, Boston, Massachusetts 02215

³SeizureTracker LLC, Lorton, Virginia 22310

⁴New York University, New York, New York 10016

⁵UCSF, San Francisco, California 94143

⁶Montefiore Medical Center/Albert Einstein College of Medicine, Bronx, New York 10467

⁷Yale University, New Haven, Connecticut 06510

⁸University of Melbourne, Fitzroy, Victoria 3065

⁹Centers for Disease Control, Washington, DC 20001

Correspondence

Daniel Goldenholz, Division of Epilepsy, Beth Israel Deaconess Medical Center, 330 Brookline Ave, Baker 5, Boston, MA 02115, USA. Tel: +1 (617) 632 8934; Fax: +1 (617) 632 8931; E-mail: daniel.goldenholz@bidmc.harvard.edu

Funding Information

This research was funded in part by the National Institutes of Neurological Disorders and Stroke, Intramural Research Division.

Received: 22 November 2017; Revised: 30 November 2017; Accepted: 6 December 2017

Annals of Clinical and Translational Neurology 2018; 5(2): 201–207

doi: 10.1002/acn3.519

Abstract

Background: There is currently no formal method for predicting the range expected in an individual's seizure counts. Having access to such a prediction would be of benefit for developing more efficient clinical trials, but also for improving clinical care in the outpatient setting. **Methods:** Using three independently collected patient diary datasets, we explored the predictability of seizure frequency. Three independent seizure diary databases were explored: SeizureTracker ($n = 3016$), Human Epilepsy Project ($n = 93$), and NeuroVista ($n = 15$). First, the relationship between mean and standard deviation in seizure frequency was assessed. Using that relationship, a prediction for the range of possible seizure frequencies was compared with a traditional prediction scheme commonly used in clinical trials. A validation dataset was obtained from a separate data export of SeizureTracker to further verify the predictions. **Results:** A consistent mathematical relationship was observed across datasets. The logarithm of the average seizure count was linearly related to the logarithm of the standard deviation with a high correlation ($R^2 > 0.83$). The three datasets showed high predictive accuracy for this log–log relationship of 94%, compared with a predictive accuracy of 77% for a traditional prediction scheme. The independent validation set showed that the log–log predicted 94% of the correct ranges while the RR50 predicted 77%. **Conclusion:** Reliably predicting seizure frequency variability is straightforward based on knowledge of mean seizure frequency, across several datasets. With further study, this may help to increase the power of RCTs, and guide clinical practice.

Key Points

- The variance of seizure frequency is predictable across diverse data.
- The logarithm of the average seizure count is highly correlated with the logarithm of the standard deviation.
- In the future, these predictions may be built into clinical trial analysis, and into clinical practice.

Introduction

Clinical trials in epilepsy have suffered from steadily rising placebo response rates over the past several decades¹ typically ranging 4–27%² but recently reaching as high as 40%.³ This can translate into unsuccessful trials, and subsequent increasing trial development costs.⁴ Natural variability in seizure frequency is a relatively unmeasured quantity. However, it may explain a portion of the “placebo effect” in epilepsy trials.⁵ Gaps in knowledge about

this variability hampers interpretation of any randomized clinical trial (RCT) that bases the outcome on seizure frequency changes. The 50%-responder rate, the preferred outcome measure of the European Medicines Agency, was selected because it is clinically relevant. However, because of the natural variability in seizure frequencies, subjects in the placebo arm may be misidentified as “responders”. Simulations based on 1767 patient seizure diaries showed that many 50%-responders in RCTs may subsequently become nonresponders (and perhaps subsequently again become responders) due to natural variability.⁵ This suggests that using the 50%-responder rate to measure “improvement” is confounded by the noise of natural variability. The signal-to-noise ratio of improvement versus natural variability is likely to be lower than desired for cost-effective RCT implementation. Understanding the expected variability in seizure rates would be of great value in improving interpretability, generalizability, and efficiency of epilepsy RCTs.

Standard clinical practice requires an implicit judgment about natural variability. Specifically, physicians make medication regimen changes based on whether a patient’s typical seizure rate has worsened above an expected upper bound on seizure frequency (decided based on clinical experience). Moreover, if a new drug adjustment results in seizure reduction below an expected rate (again decided based on clinical experience), the adjustment is considered beneficial. Therefore, any perceived drug effects are based on an implicit accounting for natural variability. For patients that achieve long-term seizure-freedom, such calculations are unnecessary. But if the seizure-freedom is short-lived, measured over a short duration, or in the absence of seizure-freedom, these calculations are currently left to the intuitive decision-making of individual practitioners, as no formal clinical tools exists.

Clinicians and trialists may benefit from a robust method for measuring/predicting the extent of the seizure rates based on natural variability. This study represents the first attempt to predict the variance of seizure frequency measurements, using a multi-modal data-driven approach.

Methods

Data

The data came from three independently collected patient diary databases (Table 1). Each dataset was managed in deidentified format, consistent with the recommendations of the NIH Office of Human Subject Research Protections, Protocol #12301. For each dataset, the data were redacted into diary format. The patients were not required to have fixed, unchanging medication regimens. In fact, some patients changed their medications often, while others did not. In the case of SeizureTracker, the medication change data was sufficiently incomplete that it was not evaluated. These diverse data provided a robust basis for our investigation into seizure rate variability, and provided confidence in the generalizability of results.

Data were obtained from a study (NeuroVista) in which subdural electrodes were chronically implanted in an attempt to provide patients with a seizure warning system.⁶ Although only 15 patients were enrolled in that study, it represents one of the most completely characterized longitudinal seizure datasets available. All patients were adults with confirmed focal epilepsy. The data consisted of several types of seizures: type 1, which were clinical seizures (reported or confirmed to be clinical by audio review) that had electrographic correlation; type 2, unconfirmed clinical (unreported) seizures with electrographic pattern identical to type 1; and type 3, subclinical, nonreported seizures with electrographic patterns that differed from types 1 and 2. Patients maintained implants for 7–24 months (median 12).

A second dataset was obtained from the Human Epilepsy Project (HEP),⁷ which is an ongoing multicenter study based on a highly screened set of adult patients with focal epilepsy who are enrolled early in their diagnosis, and has very complete data recording including self-reported data quality measures. Data included all 263 patients from July 2012 to March 2016. This dataset represents the one of the most reliable patient-reported

Table 1. Data sets. Shown here are the three datasets used for testing Model V and Model F (NeuroVista, HEP, and SeizureTracker), as well as the additional dataset (denoted with *) from SeizureTracker used in the validation simulation.

	N	N (after exclusions)	Study duration in months (median)	Diary durations after exclusion criteria	Ages	Epilepsy
NeuroVista	15	15	7–24 (12)	7–24 (12)	Adults	Focal
Human Epilepsy Project	263	93	1–46 (16)	8–42 (22)	Adults	Focal
SeizureTracker.com	12946	3016	0–596 (1)	6–596 (20)	Adults + children	Focal and generalized
SeizureTracker.com (*)	1835	403	0–8 (3)	6–8 (8)	Adults + children	Focal and generalized

seizure databases available, because of the extensive physician oversight and independent verification of diagnosis and data quality. Diary data for each patient tracked 1–46 months of data (median 16).

A third dataset was obtained from SeizureTracker.com,⁸ an online and mobile free service, representing one of the world's largest patient managed seizure diary databases. Of note, the patients in this database have focal or generalized epilepsy, and include adults and children. The SeizureTracker database consisted of a data export of all consecutive data entered from the project start in December 2007 through October 2015, comprising 12,946 patients and 1,060,680 seizures. A second export of SeizureTracker from October 2015 through May 2016 was obtained for a validation stage see Section 2.5 below adding 149,356 new seizures from 1835 patients (846 of which were new patients).

Preprocessing

To compute longitudinal predictions of seizure frequency (Mean and variance of seizure frequency), some preprocessing was required to ensure that there was both sufficient data to study, and an adequate signal-to-noise ratio. In all three datasets, we required each patient to have at least 6 months of diary data and, independently, at least six seizures recorded to be included for further analysis (see Table 1). This minimum duration was selected based on the fact that simulations were standardized to 6 months in duration.

The SeizureTracker data required additional preprocessing to reduce noise, as there was no physician curating the data. Repeated patient profiles were removed. Patients with unreported or impossible ages were excluded. Seizures reported to occur after the export dates were excluded. Seizures reported with identical start times were removed except for the first one, under the assumption that these represented erroneous repeat entries. Seizures erroneously reported to occur prior to patients' date of birth were excluded.

Mean and variance of seizure frequency

We explored the 2-week seizure frequencies of individual patients. Because seizures are very rapid events typically lasting less than two minutes,⁹ truncation of events at the edges of 2-week segments was considered unnecessary. All available 2-week segments were included in these calculations. Thus, for the j th patient, the mean (μ_j) and standard deviation (σ_j) were computed across all M available 2-week segments. For the i th segment in the j th patient, the 2-week seizure count was given by C_{ij} :

$$\mu_j = \frac{1}{M} \sum_{i=1}^M C_{ij} \quad (1)$$

$$\sigma_j = \sqrt{\frac{1}{M-1} \sum_{i=1}^M (C_{ij} - \mu_j)^2} \quad (2)$$

To account for the wide range of seizure rates, we applied a base-10 logarithmic transform to both μ_j and σ_j .

$$y_j = \log_{10}(\sigma_j) \quad (3)$$

$$x_j = \log_{10}(\mu_j) \quad (4)$$

We plotted the transformed mean versus the transformed standard deviation for each patient's seizure rates. A linear regression line was fit through the set of all patients from each dataset, with 95%-confidence regions as well. The coefficient of determination (R^2) and coefficient estimates (m and b) were reported for each dataset:

$$y_j = mx_j + b \quad (5)$$

Predicting seizure counts using the log–log plot, the “L relationship”

With a given average seizure frequency (μ_j), one could use Equations 1–5 to predict the standard deviation (σ) with the “L relationship”:

$$\sigma_j = 10^{m \log_{10}(\mu_j) + b} \quad (6)$$

To test the accuracy of such predictions, we divided each patient's diary into 6-month segments to represent a typical clinical trial of 2-months baseline, 1-month titration, and 3-months experimental period. For each 6-month segment, the 2-month “baseline” was used to estimate μ_j with Equation 1.

Two approaches for seizure frequency range predictions were tested on the individual patient level: the 50%-responder rate (RR50) method and the L method. For the L method, the 95%-confidence limits of expected experimental C_{ij} rates were computed using measured μ and Equation 6 predicted σ :

$$C_{ij} \in [\mu_j - 2\sigma_j, \mu_j + 2\sigma_j] \quad (7)$$

The RR50 model has been required by the EMA for traditional epilepsy RCTs, and therefore has been employed for many years. It makes no assumptions about the distribution (unlike the Gaussian assumption of the L model). Rather, it only specifies the lower limit of the 2-

week counts during the experimental phase from the j th patient (C_{ij}) as follows:

$$C_{ij} \in [0.5\mu_j, \infty) \quad (8)$$

Results

Of the available 12,949 patients in the first SeizureTracker export with at least one seizure recorded, preprocessing decreased this to 11,736. Then 5938 remained after requiring six or more seizures, and 3016 patients were retained after requiring 6 months or longer diaries. Of the 263 patients from HEP, 107 had six or more seizures, and of those, 93 had 6 months or longer recorded. All 15 NeuroVista patients were included.

The plots relating average 2-week seizure frequency to standard deviation (the square root of variance) from both patient-reported datasets (SeizureTracker and HEP) and the confirmed clinical (type 1) seizures from intracranial recordings (NeuroVista) superimposed are shown in Figure 1A, with linear fit lines. The three forms of NeuroVista seizures are shown in Figure 1B. Note that each point in Figure 1 represents the entire diary of an individual. Although coming from very different sources, all the log-transformed data followed a consistent linear relationship with a high coefficient of determination (0.827–0.971). Notably, all the fit lines overlapped, suggesting a common trend across otherwise disparate datasets. Using the results of Figure 1, we selected $m = 0.7$ and $b = 0.0097$ for Equations 5 and 6.

The validation data export of SeizureTracker included 1835 patients, of which 403 patients met inclusion criteria. A plot of the prediction accuracy of the ranges from the two methods is shown in Figure 2. The accuracy of the log–log relation ranged from 96–100%, while the RR50 method yielded 42–70% accuracy. In the validation dataset, the log–log predicted 94% of the correct ranges while the RR50 predicted 77%.

Discussion

Our study found evidence that changes in seizure frequency could be accounted for with high accuracy using a Gaussian model coupled to a non-linear predictor. The predictor was generated based on a consistent relationship noted across three independent datasets. Future work could lead to broader implications: smaller and less expensive clinical trials, and improved clinical care models.¹⁰

Advantages of variance prediction

Current epilepsy RCTs assume that any reduction in baseline seizure rate below 50% represents an improvement. This implies a linear relationship between the expected range and average seizure frequency (Equation 8). Our investigation of three independent datasets (Fig. 1) found that a non-linear relationship (Equations 1–7) is more appropriate. Correcting for these expected levels of variability may increase statistical

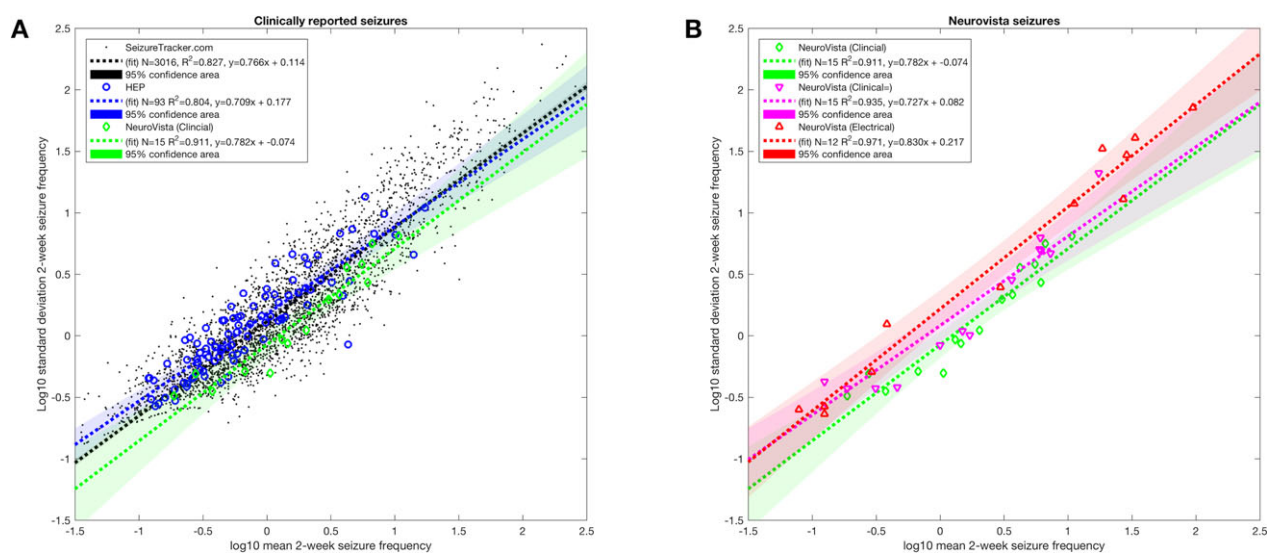


Figure 1. Log–log plot of seizure frequency mean and standard deviation (the square root of variance) for each patient. Each patient is represented by a single point on this plot. Linear fit lines (with confidence regions) are drawn for each of the datasets. A: Representative datasets: the clinically reported and verified seizures (subtype 1) of NeuroVista, the HEP data, and the SeizureTracker data. B: The three subtypes of NeuroVista, plotted in the same way as A. These plots were used to develop the predictions in Equation 6.

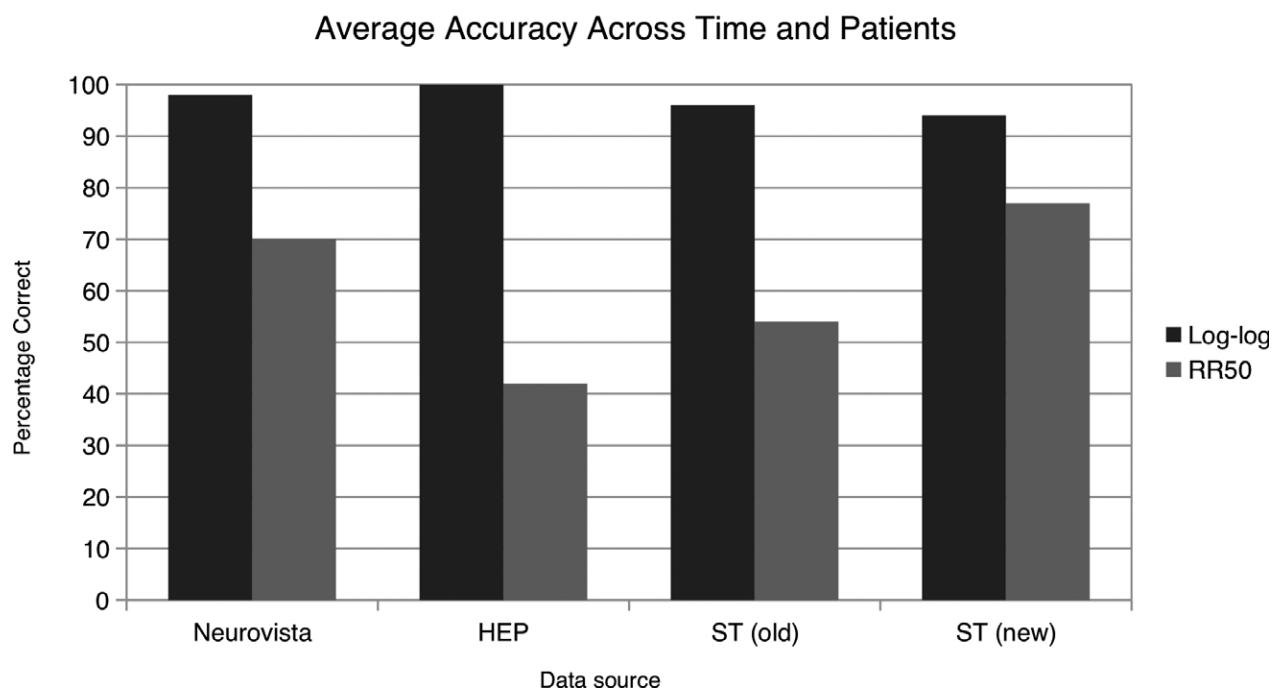


Figure 2. Predictions from the 50%-responder (RR50) method (Equation 8) and log–log method (Equation 7), applied to multiple datasets to estimate the range of possible seizure frequencies. If a seizure frequency was within the predicted range, then it was scored as correct. ST (old) is the large SeizureTracker dataset used in Figure 1. ST (new) is the independent validation dataset not included in the exploratory analysis from Figure 1. Broadly speaking, the log–log predictions had considerably more accuracy than the RR50 predictions across all datasets assessed.

efficiency of clinical trials, thereby reducing costs. One example use would be to integrate these variability predictions as an extension to the therapeutic intensive seizure analysis method (TISA), allowing for more nuanced statistical analysis with small numbers of patients, brief trial duration and objective seizure recording methods.^{11,12} More broadly, these variability corrections could be built into a trial analysis calculation, similar to the recently reported Z_V method.¹⁰

In the outpatient clinic, physicians treating patients with different baseline seizure frequencies can quantitatively anticipate natural fluctuations in seizure frequency, and thereby evaluate treatment with greater confidence. An explicit account of expected variance will ensure patients and clinicians do not respond precipitously to “changes” in seizure frequency. In addition, the overwhelming consistency of the log–linear relationship across patients indicates that a single estimate of the patient’s seizure–frequency is all that is required to obtain the predicted variance of their seizure rates. Therefore, very little baseline data is required to apply our model for seizure variability on a patient-specific basis. Furthermore, the model L calculation can be computed easily in the clinic, on a mobile app, or within modern electronic medical record systems. Although not currently available clinically, future work will explore this possibility.

Limitations of the prediction model

The model has some hidden assumptions that require consideration. First, it assumes that the relationship between seizure frequency and variance is predictable. Despite the reproducibility across datasets and the validation dataset, the possibility remains that these findings will not fully generalize across all forms of epilepsy and in all circumstances. Second, it assumes that the baseline measurement is sufficient to obtain an estimate of the true average seizure frequency. This estimate is imprecise; therefore, the question of how much uncertainty in estimated seizure frequency can be tolerated by the model should be a topic for future investigation. The fact that predictions of future seizure variability were accurate (Fig. 2), indicate that the degree of precision in the baseline estimate may be sufficient.

An additional limitation related to the nature of the data we evaluated should be considered. Given that we required patients to have six or more seizures recorded during six or more months, we were excluding those patients that have very few lifetime seizures. Moreover, it is well-recognized that datasets that describe very infrequent seizures are not currently available. As a result of these concerns, it is possible that the log–log relationship described in this study may not relate to patients with

very few lifetime seizures or very infrequent seizures. The results are expected to be biased toward patients more likely to be drug resistant. As well treated patients change from a state of having many seizures to extremely few or none, it is unknown if the log–log relationship would be meaningful after that transition.

Data considerations

One of the most compelling points of the results we have presented was their consistency, despite the fact that the three datasets used in this study were from diverse sources. The NeuroVista data was derived from very few patients with medication resistant focal epilepsy, whereas the other two datasets included focal and generalized forms of epilepsy. One of the key strengths of NeuroVista is that the data can be considered gold standard in terms of reliability of seizure detection, since intracranial electrodes were used to identify and characterize each seizure. The HEP dataset requires that patients enroll early in their diagnosis of epilepsy, whereas the other two do not. HEP data was composed of patient reported outcomes. Nevertheless, it had the most detailed mechanisms in place for multiple physicians reviewing clinical data, ensuring reliability. The SeizureTracker dataset includes longitudinal data that spans many years and more patients than most existing datasets in the world, while the other two are more restricted. The SeizureTracker dataset is the only one of the three that does not have physician oversight to ensure data reliability. Despite these various differences, a number of common results emerged, which strengthen the claims that these findings are generalizable.

Unlike HEP and NeuroVista, SeizureTracker data has additional biases inherent in any self-reported patient database lacking physician oversight.⁸ Perhaps the most challenging is “diary fatigue,” that is, the gradual or abrupt disuse of the diary because the patient or caregiver loses interest. There is no straightforward correction available. SeizureTracker also was unique among the three datasets because of the population studied. That data uniquely included children and generalized forms of epilepsy, neither of which were included in the other two datasets. To overcome this and other biases, we have studied HEP and NeuroVista that both had considerable physician oversight, and found the results were consistent across each.

An important consideration, particularly relevant to the HEP dataset, though at least partially relevant to all of them, is the possibility of medication changes influencing seizure frequencies and variability. HEP is unique here: because these patients were recently diagnosed with epilepsy, they would be expected to have more frequent medication changes than some other populations.

Although this effect may certainly influence the outcome of the predictions, adjusting for this would be expected to only improve the estimations further. Thus, unadjusted values are presented here as a lower bound for the possibility of prediction.

It is clinically challenging to determine if a patient has failed a treatment. For instance, after taking drug X for 3 months, a patient that had a single breakthrough seizure may reasonably ask their doctor, “should I stop this drug? It doesn’t seem to work.” The methodology described in this study will not clearly answer questions of this variety, because seizure-freedom is not sufficiently sampled with the data considered here. By the same token, patients who have extended seizure-free periods for many months with a single breakthrough seizure would not benefit from the present analysis, as no meaningful prediction could be made for them. Conversely, patients that are drug-resistant that have relative decreases or increases in their “usual” seizure frequency may benefit from this type of analysis because it may allow for a more structured approach to determining what is, and what is not, a change. As more comprehensive datasets become available, the breadth with which this type of analysis would apply could expand.

Conclusions

This study represents the first formal attempt to quantify the relationship between average seizure frequency and variability. The findings presented here suggest that the new L technique has the potential to improve the power and efficiency of RCTs. Further investigation is required to validate this possibility. In the future, this could improve the safety of patients via decreased exposure to nontherapeutic doses of medications.¹³ Indeed, smaller, more efficient trials could lead to much lower drug trial costs, thereby accelerating drug discovery.

Acknowledgments

Primary data were obtained by the Human Epilepsy Project team, the NeuroVista team, and the SeizureTracker.com team. Use of the data was facilitated by the International Seizure Diary Consortium (<https://sites.google.com/site/isdchome/>).

Author Contributions

Conception and design of study: DG, WT. Acquisition and analysis of data: DG, RM, JF, DL, RK, SH, SC, KD, JH, PK, MK, AS. Drafting substantial portion of manuscript: DG, SG, WH.

Conflicts of Interest

None declared.

References

1. Rheims S, Perucca E, Cucherat M, Ryvlin P. Factors determining response to antiepileptic drugs in randomized controlled trials. A systematic review and meta-analysis. *Epilepsia* 2011;52:219–233.
2. Goldenholz DM, Goldenholz SR. Response to placebo in clinical epilepsy trials—Old ideas and new insights. *Epilepsy Res* 2016;122:15–25.
3. French JA, Krauss GL, Wechsler RT, et al. Perampanel for tonic-clonic seizures in idiopathic generalized epilepsy. A randomized trial. *Neurology* 2015;85:950–957.
4. PhRMA. Profile Biopharmaceutical Research Industry [Internet]. 2015. Available from: http://www.phrma.org/sites/default/files/pdf/2015_phrma_profile.pdf
5. Goldenholz DM, Moss R, Scott J, et al. Confusing placebo effect with natural history in epilepsy: a big data approach. *Ann Neurol* 2015;78:329–336.
6. Cook MJ, O'Brien TJ, Berkovic SF, et al. Prediction of seizure likelihood with a long-term, implanted seizure advisory system in patients with drug-resistant epilepsy: a first-in-man study. *Lancet Neurol* 2013;12:563–571.
7. French J, Kuzniecky R, Lowenstein D. The Human Epilepsy Project [Internet]. [date unknown]; Available from: <http://www.humanepilepsyproject.org/>
8. Fisher RS, Blum DE, DiVentura B, et al. Seizure diaries for clinical research and practice: limitations and future prospects. *Epilepsy Behav* 2012;24:304–310.
9. Theodore WH, Porter RJ, Albert P, et al. The secondarily generalized tonic-clonic seizure: a videotape analysis. *Neurology* 1994;44:1403–1407.
10. Goldenholz DM, Goldenholz SR, Moss R, et al. Does accounting for seizure frequency variability increase clinical trial power? *Epilepsy Res* 2017;137:145–151.
11. Stefan H, Wang-Tilz Y, Pauli E, et al. Onset of action of levetiracetam: a RCT trial using therapeutic intensive seizure analysis (TISA). *Epilepsia* 2006;47:516–522.
12. Stefan H, Wang Y, Pauli E, Schmidt B. A new approach in anti-epileptic drug evaluation. *Eur J Neurol* 2004;11:467–473.
13. Ryvlin P, Cucherat M, Rheims S. Risk of sudden unexpected death in epilepsy in patients given adjunctive antiepileptic treatment for refractory seizures: a meta-analysis of placebo-controlled randomised trials. *Lancet Neurol* 2011;10:961–968.