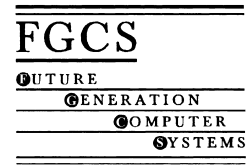




ELSEVIER

Future Generation Computer Systems 16 (2000) 351–359



Keystroke dynamics as a biometric for authentication

Fabian Monroe^{a,*}, Aviel D. Rubin^b

^a Courant Institute of Mathematical Science, New York University, New York, NY, USA

^b AT&T Labs-Research, Florham Park, NJ, USA

Accepted 3 March 1999

Abstract

More than ever before the Internet is changing computing as we know it. Global access to information and resources is becoming an integral part of nearly every aspect of our lives. Unfortunately, with this global network access comes increased chances of malicious attack and intrusion. In an effort to confront the new threats unveiled by the networking revolution of the past few years reliable, rapid, and unintrusive means for automatically recognizing the identity of individuals are now being sought. In this paper we examine an emerging non-static biometric technique that aims to identify users based on analyzing habitual rhythm patterns in the way they type. ©2000 Elsevier Science B.V. All rights reserved.

Keywords: Identity verification; User authentication; Biometrics

1. Introduction

The increasing use of automated information systems together with our pervasive use of computers has greatly simplified our lives, while making us overwhelmingly dependent on computers and digital networks. Technological achievements over the past decade have resulted in improved network services, particularly in the areas of performance, reliability, and availability, and have significantly reduced operating costs due to the more efficient utilization of these advancements. However, the overwhelming interest in global accessibility brought about by these advances in technology have unveiled new threats to computer system security. As we press into the twenty-first century, new challenges abound. Advanced safeguards against fraud and impersonation, as well as foolproof

measures against unauthorized access to computer resources and data are now being sought. We present one such safeguard based on authenticating access to computers by recognizing certain unique and habitual patterns in a user's typing rhythm.

We argue that the use of keystroke rhythm is a natural choice for computer security. This argument stems from observations that similar neuro-physiological factors that make written signatures unique, are also exhibited in a user's typing pattern [14]. When a person types, the latencies between successive keystrokes, keystroke durations, finger placement and applied pressure on the keys can be used to construct a unique signature (i.e., profile) for that individual. For well-known, regularly typed strings, such signatures can be quite consistent. Furthermore, recognition based on typing rhythm is not intrusive, making it quite applicable to computer access security as users will be typing at the keyboard anyway.

This paper presents our results for an authentication system based on the use of keystroke dynamics.

* Corresponding author.

E-mail addresses: fabian@cs.nyu.edu (F. Monroe),
rubin@research.att.com (A.D. Rubin).

Keystroke dynamics is the process of analyzing the way a user types at a terminal by monitoring the keyboard inputs thousands of times per second, and attempts to identify them based on habitual rhythm patterns in the way they type. We present our data selection and extraction methods as well as our classification and identification strategies. Our observations and findings are discussed and compared with prior work in this area.

2. Biometrics

Biometrics, the physical traits and behavioral characteristics that make each of us unique, are a natural choice for identity verification. Biometrics are excellent candidates for identity verification because unlike keys or passwords, biometrics cannot be lost, stolen, or overheard, and in the absence of physical damage they offer a potentially foolproof way of determining someone's identity. Physiological (i.e., *static*) characteristics, such as fingerprints, are good candidates for verification because they are unique across a large section of the population.

Indispensable to all biometric systems is that they recognize a *living* person and encompass both physiological and behavioral characteristics. Physiological characteristics such as fingerprints are relatively stable physical features that are unalterable without causing trauma to the individual. Behavioral traits, on the other hand, have some physiological basis, but also reflect a person's psychological makeup. Unique behavioral characteristics such as the pitch and amplitude in our voice, the way we sign our names, and even the way we type, form the basis of *non-static* biometric systems.

Biometric technologies are defined as “automated methods of verifying or recognizing the identity of a living person based on a physiological or behavioral characteristic” [19]. Biometric technologies are gaining popularity because when used in conjunction with traditional methods for authentication they provide an extra level of security. Available counter-measures to the problem of identity verification can be categorized into three main groups: those that rely on (a) something a person knows (e.g. a password), (b) something a person possesses (e.g. an ID card), or (c) characteristics of a person.

Security measures which fall under categories (a) and (b) are inadequate because possession or knowl-

edge may be compromised without discovery — the information or article may be extorted from its rightful owner. Increasingly, attention is shifting to positive identification by biometric techniques that encompass the third class of identification (i.e., biometrics) as a solution for more foolproof methods of identification. For the foreseeable future, these biometric solutions will not eliminate the need for ID cards, passwords and PINs. Rather, the use of biometric technologies will provide a significantly higher level of identification and accountability than passwords and cards alone, especially in situations where security is paramount.

2.1. Let us see your hands, eyes and face

Modern biometric schemes generally rely on aspects of the body and its behavior. Slight changes in behavior are inevitable when dealing with non-static biometrics since they are influenced by both controllable actions and unintentional psychological factors. Therefore, biometric technologies need to be robust and adaptive to change — online signature verification systems, for example, update the reference template of a user on each successful authentication to the login device to account for slight variations in the signature.

Some examples of identifying biometric features being used for identification based systems include hand geometry, thermal patterns in the face, blood vessel patterns in the retina and hand, finger and voice prints, and handwritten signatures (see [4,6,7,13,15,22,24]). Today, a few devices based on these biometric techniques are commercially available. However, some of the techniques being deployed are easy to fool, while others like iris pattern recognition, are too expensive and invasive.

In an effort to provide a passive, inexpensive and more foolproof method than traditional passwords for verifying an individual's identity, we present keystroke dynamics. The techniques presented herein rely on pattern recognition. A brief overview of fundamental concepts of pattern recognition is given in the following section.

3. Pattern recognition: representation, extraction, and classification

The design of an automatic pattern recognition system involves *representation*, *extraction*, and *classifi-*

cation. Representation of input data measures characteristics of the pattern of object to be recognized. When the measurements obtained yield information in the form of real numbers, it is often useful to think of a pattern vector as a point in an n -dimensional Euclidean space.

The *extraction* of characteristic features from the input data and the reduction of the dimensionality of the resulting pattern vectors is often referred to as the preprocessing and feature extraction problem. For example, we may choose to use only a selected number of measurements from the input, either because these features are enough to identify the individual (like the eyes and mouth) or because the addition of other extra features increase the computational complexity of the problem or yields no real benefit. The number of degrees of freedom of variation in the chosen index across the human population, their immutability over time and immunity to intervention, and the computational prospects for efficiently encoding and reliably recognizing the identifying pattern, must all be assessed during feature extraction.

Classification and identification involves the determination of optimum decision procedures. After the observed data from patterns to be recognized have been expressed in the form of measurement vectors in the pattern space, we want to decide to which pattern class these data belong [23,8]. For example, given a number of face images, and an “unknown” reference template from a database of available faces, we want to be able to positively identify the “unknown” individual with a specified level of certainty.

The analysis of personal features has a natural range of variation; a biometric method never provides an absolutely certain identification. As such, a biometric identification system can fail in one of two ways; either an authorized user is rejected or an illegitimate user can be incorrectly granted access to the system. Biometric systems must allow adjustments to control the error probabilities to some degree.

4. Keystroke dynamics: not what you type, but how you type

Keystroke dynamics is the process of analyzing the way a user types at a terminal by monitoring the keyboard inputs thousands of times per second in

an attempt to identify users based on habitual typing rhythm patterns. It has already been shown (see [14,18,20]) that keystroke rhythm is a good sign of identity. Moreover, unlike other biometric systems which may be expensive to implement, keystroke dynamics is almost free — the only hardware required is the keyboard.

The application of keystroke rhythm to computer access security is relatively new. There has been some sporadic work done in this arena. Joyce and Gupta [14] present a comprehensive literature review of work related to keystroke dynamics prior to 1990. We briefly summarize these efforts and examine the research that has been undertaken since then.

4.1. The current state of keystroke dynamics

Keystroke verification techniques can be classified as either static or continuous. Static verification approaches analyze keystroke verification characteristics only at specific times, for example, during the login sequence. Static approaches provide more robust user verification than simple passwords, but do not provide continuous security — they cannot detect a substitution of the user after the initial verification. Continuous verification, on the contrary, monitors the user’s typing behavior throughout the course of the interaction.

As early as 1980, researchers have been studying the use of habitual patterns in a users typing behavior for identification. To our knowledge, Gaines et al. [9] were the first to investigate the possibility of using keystroke timings for authentication. Experiments were conducted with a very small population of seven secretaries. A test of statistical independence of their profiles was carried out using the *T-Test* under the hypothesis that the means of the digraph times at both sessions were the same, but the variances were different. Similar experiments were conducted by Leggett et al. [16,17] with seventeen programmers but for the continuous approach to user verification. The authors report an identity verifier that validates the results of [9] — an identity verification system with false alarm rate of about 5.5 percent and impostor pass rate of approximately 55.0 percent.

While the approaches of Gaines et al. [9] and Leggett et al. [16,17] address a number of problems inherent with identity verification via keystroke

timings, there was considerable room for improvement. For example, the pool variance estimate used in [17] is meaningful only when there is homogeneity of variance across all reference digraph latencies; however studies by Mahar et al. [18] show that there is a significant variability with which typists produce each digraph, and hence the use of a pooled estimate digraph latency variability is inappropriate.

An additional limitation of the digraph latency based technique [17] is the use of a single low-pass temporal filter for all typists for the removal of outliers. The rationale for this approach is that digraphs with abnormally long latencies are not likely to be representative of the authorized user's typing. While this seems like a reasonable assumption it has recently been shown ([18,20]) that one filter value for all typists does not yield optimal performance.

Furthermore, empirical data from Gentner [10] suggests that the median interkey latency of expert typists is approximately 96 ms, while that of novice typists is near 825 ms. Therefore, the 500 ms low-pass filter used by [17] excludes many keystrokes typical of novice typists, while at the same time, includes many keystrokes which are not representative of an expert typist [18]. Studies by [18,20] showed that the use of digraph-specific measures of variability instead of one low-pass filter can lead to measurable improvements in verification accuracy. Moreover, the approach of [17] to keystroke verification uses the key down-to-down time as the base unit of measure, but this measure may be further delineated into two orthogonal components — total time the first key is depressed (i.e. keystroke duration), and the time between a key is released and the next key is pressed (i.e. keystroke latency). Previous works [3,18,20] used these two components in their verification systems. However, the initial sample sets of [3,20] did not provide enough data to ascertain whether the use of the two separate orthogonal digraph components added significant predictive power to the more traditional key down-to-down measure. Substantially improved performance results based on using the bivariate measure of latency with an appropriate distance measure were achieved by [18].

Some neural network approaches [1,3,12] have also been undertaken in the last few years. While the back-propagation models used yield favorable performance results on small databases, neural networks have a fundamental limitation in that each time a new

user is introduced into the database, the network must be retrained. For applications such as access control, the training requirements are prohibitively expensive and time-consuming. Furthermore, in situations where there is a higher turnover of users, the down time associated with retraining can be significant.

A promising research effort in applying keystroke dynamics as a static authentication method is the work of Joyce and Gupta [14]. Their approach is relatively simple and yields impressive results. Our work extends that of Joyce and Gupta and we review their classifier in Section 4.4.

4.2. Data selection and representation

The performance results reported here are based on a database of profiles collected over a period of 11 months. Data for 63 users was collected at a variety of Sun Workstations at NYU and Bell Communications Research. Typing proficiency was not a requirement in this study although almost all participants were familiar with computers. Unlike previous studies in which the observers had complete control over the collection of the data [2], participants ran the experiment from their own machines at their convenience. Participants down-loaded and executed the experiment on their local machines and the results were automatically encoded and electronically mailed back to us. Fig. 1

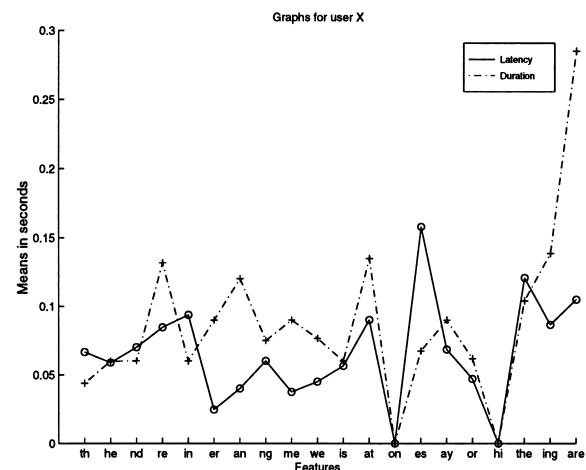


Fig. 1. Example reference profile. The top n most frequent features in the pattern vector are shown on the X-axis. The users keystroke latencies, as well as keystroke durations, are graphed above. The graphs show that on an average, the user suppresses keys for a longer period than it takes him/her to type them.

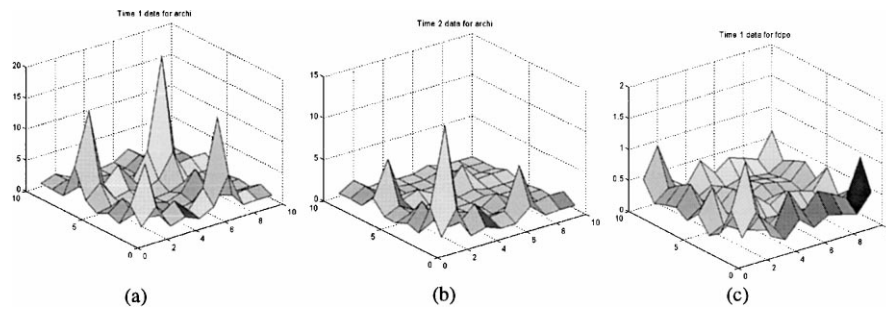


Fig. 2. Plots (a) and (b) depict the covariance matrices for the same user at two different time intervals across the same set of features. Plot (c) shows the covariance matrix for a different user over the same set of features. Notice the same peaks and structure between plots (a) and (b) which is quite distinct from the profile in (c).

shows an example of a profile received for a user in the data set. An alternate representation showing plots of the covariance matrices (of the keystroke latencies for a particular feature set) for different users over different time intervals is shown in Fig. 2.

4.3. Data extraction

To evaluate the behavior and performance of each of the classifiers presented in Section 4.4 we developed a C++ toolkit for analyzing the data. The toolkit was built using the xview library routines, and serves as a frontend to the main recognition engine. The toolkit is helpful in diagnosing system behavior and can generate graphical output for both the Matlab and Gnuplot systems. Fig. 3 is from the main panel of the toolkit.

The data extraction toolkit provides a quick way to establish rough properties on the data set by partitioning the users in distinct groups. Our clustering criterion represents a heuristic approach that is guided by intuition — users are clustered into groups comprising of (possibly) disjoint feature sets in which the features in each set are pairwise correlated.

Feature sets are determined through factor analysis (FA) [5]. Factor analysis seeks a lower dimensional representation that accounts for the correlation among features. This idea partitions the database of users into subsets whose in-class members are “similar” in typing rhythm over a particular set of features and whose cross-class members are dissimilar in the corresponding sense. For example, members of group i may exhibit strong individualistic typing patterns for features in the set $S = \{\text{th, ate, st, ion}\}$, whereas members of group j may be more distinctive over the features

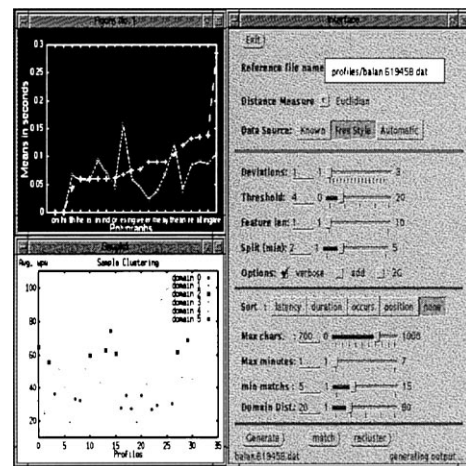


Fig. 3. To automate the data selection and extraction process a system toolkit was designed to assist in the visualization, tuning, and overall analysis of the data. A graphical user interface with various tunable options allow the operator to diagnose the performance of each of the classifiers in detail. The above is a snapshot from the main panel of the interface.

$S = \{\text{ere, on, wy}\}$. K-Nearest Neighbor [8] is used as the clustering algorithm. The net result is a hierarchical cluster that assists in user identification.

4.4. Classification and identification

The problem of recognizing a given pattern as belonging to a particular person either after exhaustive search through a large database, or by simply comparing the pattern with a single authentication template can be formulated within the framework of statistical decision theory. By this approach one can convert the

problem of pattern recognition into a much more expedient task, which involves the execution of tests of statistical independence. The approaches described in the following paragraphs adhere to this model.

The classification technique employed by Joyce and Gupta [14] represents the mean reference signature for a given user as $\mathcal{M} = \{\mathcal{M}_{\text{username}}, \mathcal{M}_{\text{password}}, \mathcal{M}_{\text{firstname}}, \mathcal{M}_{\text{lastname}}\}$. Verification is performed by comparing the test signature $\mathcal{T}$ (acquired at login time) with $\mathcal{M}$ and determining the magnitude of difference between the two profiles. Given $\mathcal{M} = (m_1, m_2, \dots, m_n)$ and $\mathcal{T} = (t_1, t_2, \dots, t_n)$ where n is the total number of latencies in the signature, the verifier computes the magnitude of difference using an L_1 norm. Positive identification is declared when this difference is within a threshold variability of the reference signature. The mean and standard deviation of the norms $\|M - S_i\|$, where S_i is one of the eight training signatures, are used to decide the threshold for an acceptable difference vector between a given $\mathcal{T}$ and $\mathcal{M}$.

Although these absolute verification rates are encouraging, Joyce and Gupta tested using a replacement methodology, which means that the distribution of the training set is necessarily representative of the learning set. The use of separate data sets, recorded at different times, would be more reliable. Therefore, we investigated the performance of classifiers based on studies where users were allowed to participate in experiments conducted at varied times under no supervision. The reference profiles collected were represented as N -dimensional feature vectors and processed in a manner similar to that of [14]. The data was split into learning and testing sets. Then, the following classifiers were used for recognition.

- Euclidean distance measure: “similarity” is based on the Euclidean distance between the pattern vectors. Let $R = [r_1, r_2, \dots, r_N]$ and $U = [u_1, u_2, \dots, u_N]$ then the Euclidean distance between the two N -dimensional vectors U and R , is defined as:

$$D(R, U) = \left[\sum_{i=1}^N (r_i - u_i)^2 \right]^{1/2}.$$

For an “unknown” U (i.e., from the testing set) the pairwise Euclidean distances $D(R_i, U)$, $i = 1, 2, \dots, n$, where n = number of pattern vectors

in the database, that were rank ordered and the profile with the minimum distance to U was chosen.

- Non-weighted probability: let U and R be N -dimensional pattern vectors as defined previously. Furthermore, let each component of the pattern vectors be the quadruple $\langle \mu_i, \sigma_i, o_i, X_i \rangle$, representing the mean, standard deviation, number of occurrences, and data value for the i th feature. Assuming that each feature for a user is distributed according to a normal distribution, we calculate the score between a reference profile R and unknown profile U as:

$$\text{Score}(R, U) = \sum_{i=1}^N S_{u_i}$$

where,

$$S_{u_i} = \frac{1}{o_{u_i}} \left[\sum_{j=1}^{o_{u_i}} \text{Prob} \left(\frac{X_{ij}^{(u)} - \mu_{r_i}}{\sigma_{r_i}} \right) \right]$$

and $X_{ij}^{(u)}$ is the j th occurrence of the i th feature of U .

In other words, the score for each u_i is based on the probability of observing the value u_{ij} in the reference profile R , given the mean (μ_{r_i}) and standard deviation (σ_{r_i}) for that feature in R . Intuitively we assign higher probabilities to values of u_i that are close to μ_{r_i} and lower probabilities to those further away. The “unknown” vector is then associated with the *nearest neighbor* in the database, i.e., to the person who maximizes the probability of the feature vector.

- Weighted probability measure: some features are more reliable than others simply because they come from a larger sample set or have a relatively higher frequency in the written language; example in English *er*, *th*, *re* should constitute greater weights than *qu* or *ts*. Thus, the notion of weights was incorporated, and the score between profiles R and U was computed as:

$$\text{Score}(R, U) = \sum_{i=1}^N (S_{u_i} * w_{u_i}),$$

where the weight of the feature u_i is the ratio of its occurrences relative to all other features in the pattern vector U . Features that are based on

many occurrences are considered more reliable and weighted higher than those features that come for a smaller sample set. Assuming that each feature for a user is distributed according to a normal distribution, a likelihood score between a reference profile R and unknown profile U is calculated based on the probability of observing a feature value in the reference profile R , given the mean and standard deviation for that feature in R . Scores are weighted and the “unknown” profile is then associated with the *nearest neighbor* in the database, i.e., the person who maximizes the score of the feature vector.

The correct identification rate using the weighted probabilistic classifier was approximately 87.18% on a dataset of 63 users¹, which represents improvement with respect to the performance of the Euclidean distance (83.22%) and the non-weighted scoring approach (85.63%). Additionally, our research argues in favor of the use of structured text instead of allowing users to type arbitrary text (i.e., text-independent or “free-text”). While recognition based on free-text may be more desirable, free-text recognition did not perform as well as recognition based on fixed-text. Recognition based on free-text may be expected to vary greatly under operational conditions in which the user may be absorbed in a task or involved in an emotionally charged situation. The fact that the input is unconstrained, that the user may be uncooperative, and that environmental parameters are uncontrolled impose limitations on what can be achieved with free-text recognition.

The superior performance of Bayesian-like classifiers for a variety of recognition tasks lead to the implementation of a Bayesian-like classifier. The approach aims to characterize the performance of the feature-based technique as a function of the number of classes to be discriminated. We assume that the feature vectors are distributed according to a Gaussian distribution and an unknown vector is associated with the person who *maximizes* the probability of the measurement vector. The classifier is defined as follows:

- let x_i be the feature vector, σ_i the interclass dispersion vector and w_i the weight vector, then the dis-

tance of two feature vectors x_i and x'_i is expressed as:

$$\Delta^\alpha(x, x') = \sum_{i=1}^n w_i \left(\frac{|x_i - x'_i|}{\sigma_i} \right)^\alpha.$$

The feature vectors, $x_1, \dots, x_n$, are derived from the sets computed by FA (see Section 4.3). In accordance with Huber [11] the value of α can be adjusted to achieve more robustness — the net effect is a slight improvement in recognition for values of α close to 1 rather than 2 as justified by the Gaussian assumption. The correct identification performance using the Bayesian classifier was approximately 92.14%, representing an improvement of almost 5% over the weighted classifier.

While it is difficult to give a meaningful comparison of our approach with that of [14,17,18] as there is no unified data set under which the approaches can be compared, overall, our results validate that of previous research and suggest that it is possible to use keystroke dynamics to accurately verify computer users, albeit in somewhat of a controlled environment.

5. Applications

Keystroke dynamics has many applications in the computer security arena. One area where the use of a static approach to keystroke dynamics may be particularly appealing is in restricting root level access to the master server hosting a Kerberos [21] key database. Any user accessing the server is prompted to type a few words of a pass phrase in conjunction with his/her username and password. Access is granted if his/her typing pattern matches within a reasonable threshold of the claimed identity. This safeguard is effective as there is usually no remote access allowed to the server, and the only entry point is via console login.

Alternatively, dynamic or continuous monitoring of the interaction of users while accessing highly restricted documents or executing tasks in environments where the user must be “alert” at all times (for example air traffic control), is an ideal scenario for the application of a keystroke authentication system. Keystroke dynamics may be used to detect uncharacteristic typing rhythm (brought on by drowsiness, fatigue etc.) in the user and notify third parties.

¹ The results reported here reflect a larger sample set than that use in [20].

6. Summary

In this paper we address the practical importance of using keystroke dynamics as a biometric for authenticating access to workstations. Keystroke dynamics is the process of analyzing the way users type by monitoring keyboard inputs and authenticating them based on habitual patterns in their typing rhythm. We review the current state of keystroke dynamics and present classification techniques based on template matching and Bayesian likelihood models.

We argue that although the use of a behavioral trait (rather than a physiological characteristic) as a sign of identity has inherent limitations, when implemented in conjunction with traditional schemes, keystroke dynamics allows for the design of more robust authentication systems than traditional password based alternatives alone. The inherent limitations that arise with the use of keystroke dynamics as an authentication mechanism are attributed to the nature of the reference “signature” and its relationship to the user — recognizing users based on habitual rhythm in their typing pattern uses dynamic performance features that depend upon an act — the rhythm is a function of the user and the environment.

The problem with keystroke recognition is that unlike non-static biometrics (such as voice) there are no known features or feature transformations which are *dedicated solely* to carrying discriminating information. Fortunately, in the past few years researchers [14,18,20] have presented empirical findings that show that different individuals exhibit characteristics in their typical rhythm that are strikingly individualistic and that these characteristics can be successfully exploited and used for identification purposes.

The performance of our classifiers on a dataset of 63 users ranges from 83.22% to 92.14% accuracy depending on the approach being used. Our research supports the observation of Mahar et al. [18] in that there is significant variability with which typists produce digraphs. Hence, we suggest the use of digraph-specific measures of variability instead of single low-pass filters. Additionally, we argue in favor of the use of structured text instead of allowing users to type arbitrary text (i.e., “free-text”) during the identification process. While recognition based on free-text may be more desirable, free-text recognition was observed to vary greatly under operational conditions; the fact that

the input is unconstrained, that the user may be uncooperative, and that environmental parameters that are uncontrolled impose limitations on what can be achieved with free-text recognition.

Acknowledgements

We would like to thank Angela Adotola, Arash Baratloo, Ivan Blank, Dan Boneh, Eric Briel, Aileen Chang, Juan Carlos, Po-Yu Chen, Churngwei Chu, Brian Coan, Daria Depiro, Sven Dietrich, Thomas Duzan, Ron Even, Jing Fan, Sabine French, Martin Garcia, Deepak Goyal, Amy Greenwald, Andy Hagerty, Andy Howell, Mehmet Karaul, Peter Kleinmann, Michael Lewis, Arjen Lenstra, Tyng-Luh Liu, Charles Mayor, Burette McLeod, Tyrone McGhee, Madhu Nayakkankuppam, Kizito Ofornagoro, Fritz Orneas, Kouji Ouchi, Gardner Patton, Toto Paxia, Bill Pink, Prabhala Prashant, Rick Porter, Raju Jawalekar, Errol Roberts, Archisman Rudra, Marc Schwarz, David Shallcross, Liddy Shriver, Julia Stone, Vijay Sreedhar, Marek Teichmann, Peter Wyckoff, Ze Yang, Shu-man Yang, and all the NYU Master students and Bellcore employees who participated in our study but for whom we do not have their full names as they no longer have login accounts at the respective institutions.

References

- [1] T.J. Alexandre, Biometrics on smartcards: an approach to keyboard behavioral signature, in: Second Smart Card Research and Advanced Applications Conference, 1996.
- [2] S. Bleha, C. Slivinsky, B. Hussein, Computer-access security systems using keystroke dynamics, *IEEE Trans. Pattern Anal. Mach. Intell. PAMI-12* (12) (1990) 1217–1222.
- [3] M. Brown, S.J. Rogers, User identification via keystroke characteristics of typed names using neural networks, *Int. J. Man-Mach. Stud.* 39 (6) (1993) 999–1014.
- [4] R. Chellappa, C.L. Wilson, S. Sirohey, Human and machine recognition of human face images: A survey, in: *Proceeding of the IEEE*, Vol. 83, 1995, pp. 705–741.
- [5] E. Cureton, Factor analysis, an applied approach, Erlbaum Associates, Hillsdale, NJ, 1983.
- [6] J.G. Daugman, High confidence visual recognition of persons by a test of statistical independence, *IEEE Trans. Pattern Anal. Mach. Intell.* 15 (11), (1993).
- [7] G.R. Doddington, Speaker recognition-identifying people by their voices, *Proceedings of the IEEE*, Vol. 73, issue no. 11, 1985, 1651–1664.

- [8] R. Duda, Pattern classification and scene analysis, Wiley, New York, 1973.
- [9] R. Gaines, W. Lisowski, S. Press, N. Shapiro, Authentication by keystroke timing: some preliminary results. Rand Rep. R-2560-NSF, Rand Corporation, 1980.
- [10] Gentner, Keystroke timing in transcription typing, *Cognitive Aspects of Skilled Typewriting*, 1993, pp. 95–120.
- [11] P.J. Huber, *Robust Statistics*, Wiley, New York, 1981.
- [12] B. Hussien, R. McLaren, S. Bleha, An application of fuzzy algorithms in a computer access security system, *Pattern Recog. Lett.* 9 (1989) 39–43.
- [13] D.K. Isenor, S.G. Zaky, Fingerprint identification using graph matching, *Pattern Recog.* 19 (2) (1986) 113–222.
- [14] R. Joyce, G. Gupta, Identity authorization based on keystroke latencies, *Commun. ACM* 33 (2) (1990) 168–176.
- [15] L.G. Kersta, Voiceprint identification, *Nature* 196 (1962) 1253–1257.
- [16] J. Leggett, G. Williams, Verifying identity via keystroke characteristics, *Int. J. Man-Mach. Stud.* 28 (1) (1988) 67–76.
- [17] J. Leggett, G. Williams, D. Umphress, Verification of user identity via keystroke characteristics, *Human Factors in Management Information Systems*, p. 89.
- [18] D. Mahar, R. Napier, M. Wagner, W. Lavery, R. Henderson, M. Hiron, Optimizing digraph-latency based biometric typist verification systems: inter and intra typists differences in digraph latency distributions, *Int. J. Human-Comp. Stud.* 43 (1995) 579–592.
- [19] B. Miller, Vital sings of identity, *IEEE Spectrum*, 1994, pp. 22–30.
- [20] F. Monrose, A. Rubin, Authentication via keystroke dynamics. Fourth ACM Conference on Computer and Communications Security, 1997, pp. 48–56.
- [21] C. Neuman, T. Ts'o, An authentication service for computer networks, *IEEE Commun.* 32 (9) (1994) 33–38.
- [22] F.J. Prokoski, R.B. Riedel, J.S. Coffin, Identification of individuals by means of facial thermography, 1992.
- [23] J.T. Tou, R.C. Gonzalez, *Pattern recognition principles*, Addison-Wesley, Reading, MA, 1981.
- [24] J. Zhang, Y. Yan, M. Lades, Face recognition: eigenface, elastic matching and neural nets, *Proceedings of the IEEE*, vol. 85, issue no. 9, September 1997, pp. 1423–1435.



Fabian Monrose is currently a Ph.D. candidate in computer science at New York University. His research interests include distributed computing, cryptography and network security. He is particularly interested in protocols that support strong mutual authentication in untrusted environments, electronic commerce, and identification technologies that rely on biometrics.



Aviel D. Rubin is a Senior Technical Staff Member at AT&T Labs, Research in the secure systems research department, and an Adjunct Professor of Computer Science at New York University, where he teaches cryptography and computer security. He is the co-author of the *Web Security Sourcebook*. Aviel holds a B.S., M.S.E., and Ph.D. from the University of Michigan in Ann Arbor (1989, 1991, 1994) in Computer

Science and Engineering. He has served on several program committees for major security conferences and as the program chair for USENIX Security '98, USENIX Technical '99, and ISOC NDSS 2000.