

Success Rates Doubled with Only One Character: Mask Password Guessing

Yunkai Zou, Ding Wang, Fei Duan

College of Cryptology and Cyber Science, Nankai University, Tianjin 300350, China; wangding@nankai.edu.cn
Key Laboratory of Data and Intelligent System Security (NKU), Ministry of Education, Tianjin 300350, China
Tianjin Key Laboratory of Network and Data Security Technology, Nankai University, Tianjin 300350, China

Abstract—While traditional whole-password guessing attacks have been extensively studied, few studies have explored mask password guessing, where an attacker has somehow obtained *partial information* of the target victim’s password (e.g., length and/or some characters) by exploiting various side-channel attacks (e.g., shoulder surfing, smudge, and keystroke audio feedback).

To evaluate the threats posed by mask attackers with varied capabilities, we investigate *four major mask guessing scenarios*, each of which is based on different kinds of information exploited by the attacker (e.g., the length of the victim’s password and some characters). For the first time, we systematically and comprehensively characterize the impacts of mask guessing that incorporate side-channel priors, personally identifiable information (PII), and previously leaked (sister) passwords, by proposing two password models: neural network-based PassSeq and probability statistics-based Kneser-Ney. Using the maximum likelihood estimation technique, we propose a new guess number estimation method to accurately and efficiently estimate the guess number required against the target password under a given password model. Extensive experiments on 15 large-scale datasets demonstrate the effectiveness of PassSeq and Kneser-Ney. Particularly, within ten guesses: (1) When a trawling attacker knows the character composition (without order) of the victim’s 4-digit PIN, the success rate increases by 152% (from 14% to 35%); (2) When a PII-based targeted attacker knows the length of the victim’s password, the success rate increases by 47%-82%; and (3) if this targeted attacker further knows *one character* of the victim’s password (besides the length), the success rate generally *doubles*, reaching 7%-29% (and these figures will be 33%-73% for a targeted attacker that can exploit the victim’s sister password).

To further validate the practicality of our mask guessing models, we collect real-world keystroke audio data from 11 popular keyboards (e.g., Apple, Dell, Lenovo) and replicate attacks where partial password information is inferred via acoustic side channels. Experiments show that our PassSeq significantly boosts the success rates of existing keystroke inference attacks, achieving an *additional 5.6%-166.7% improvement* within 10 guesses. This work highlights that mask password guessing is a damaging threat that deserves more attention.

I. INTRODUCTION

Text passwords remain the most widely used authentication method due to their simplicity of use, ease of change, and low deployment cost, and they are likely to maintain the

dominant position in the foreseeable future [12], [13], [28], [75]. Accurately modeling password guessability can not only help precisely evaluate the security threats brought by attackers, but also help administrators design and deploy effective security mechanisms (e.g., password strength meters [15], [68] and honeywords [35], [65]) to protect systems and users. A common way to evaluate the strength of a password is to estimate its guess number under a given password guessing model/algorithm/tool [18], [37].

There have been dozens of password models designed for whole-password guessing, such as probabilistic context-free grammar (PCFG [67]), Markov [43], [48] and recurrent neural network (RNN)-based models [46], [51]. However, these models are *not* optimal for mask guessing, where the attacker has further obtained partial information of the victim’s password (e.g., some characters and/or length). In general, password models like RFGuess [66], Markov [43], and RNN [46] are trained to predict the next character based on preceding characters in a password (i.e., characters are trained and predicted from a single direction). As a result, they are unable to utilize characters that follow a template to predict preceding masked characters. For example, in the template ***veu4eve** (corresponding to the password *loveu4ever*, a stylized form of ‘love you forever’), where *** represents masked characters, they cannot utilize the following *veu4eve* to determine the probabilities of the first two masks.

Worse still, existing side-channel attacks (e.g., shoulder surfing and smudge attacks) pose significant threats by potentially exposing sensitive parts of the target victim’s password. Among these attacks, shoulder-surfing attacks [11], [22], [39] are particularly concerning, as they enable attackers to directly see partial characters and/or the length of the target victim’s password. In 2024, Hu et al. [33] systematically investigated password-masking practices across modern authentication systems and found that all evaluated mobile browsers implement *dynamic masking* (i.e., briefly revealing the most recently typed character before re-masking it). This mechanism makes certain characters trivially observable to nearby shoulder-surfers or camera-based side channels. Such momentary exposure on mobile devices directly motivates our mask guessing scenarios, in which attackers exploit short-lived character leaks to substantially increase guessing success rates.

Actually, shoulder-surfing attacks are frequently encountered in real-world scenarios, especially in public or semi-



Fig. 1. Access control keypads that leave/leak input traces in the real world. We can clearly see that the left user’s password consists of the four digits 1, 5, 6, and 7, while the right user’s password contains 2, 6, and 7. Using this side-channel information, the attacker can perform mask password guessing to significantly increase her guessing success rate (see Fig. 5).

public spaces, and have alarmingly high success rates. For instance, Honan [31] reports that 85% of participants admitted being able to observe sensitive information on others’ computer screens; the Ponemon Institute [1] finds that 91% of 157 visual hacking attempts were successful in office settings. Large-scale user studies further corroborate these findings: Marques et al. [45] report that about 20% of U.S. adults (and up to 52% of young users) admit to successfully snooping on others’ phones within a single year. Aviv et al. [8] establish empirical baselines for mobile shoulder-surfing, showing success rates of 10.8% for 6-digit PINs after a single observation and 26.5% after multiple observations.

Besides shoulder surfing, there are also common and practical threats that motivate our work. For example, smudge attacks [8] allow attackers to determine the character composition of a victim’s password by analyzing the traces left/leaked on input devices (see Fig. 1); keystroke audio feedback [9], [10], [14], [74] enables attackers to listen to keystrokes, revealing the probability distribution of individual characters and the password length; Thermal attacks [4], [6], exploiting keyboard heat maps, can reliably identify the last input password character with certainty. Notably, Alotaibi et al. [6] illustrate, through an example decoding path (Fig. 11), that the per-state decoding confidence along the inferred path can follow a monotonic increasing trend with position, with later positions showing higher conditional probabilities; Bicycle attacks [24], which monitor encrypted traffic between users and authentication servers, can infer the exact password length. Particularly, Harsha et al. [24] reveal that >84% of the Alexa Top-100 pages are vulnerable to password-length leakage.

If the attacker knows some characters and the exact length of the victim’s password *with certainty*, she can directly perform template-based mask guessing. Alternatively, if the attacker has partial information about some characters (e.g., the probability distribution of characters obtained through keystrokes or the composition of characters obtained through smudges), she can incorporate this distribution (derived from acoustic side-channel analysis) as prior knowledge into the

mask-based model to further improve the guessing success rate based solely on sound. In Sec. IV-F, we demonstrate that *our proposed password model, when integrated with acoustic-derived priors, can significantly boost the guessing success rates of existing acoustic side-channel attacks (e.g., [3], [23]) by 6%-167% within 10 guesses (see Table IX)*. This is achieved using keystroke data collected from 11 real-world keyboards (e.g., Apple and Dell laptops).

Despite the increasing deployment of multi-factor authentication (MFA), password-only authentication remains prevalent. Recent reports indicate that merely 27%-34% of small and medium-sized enterprises have adopted MFA [57], and that about 60% of users still rely on passwords for their personal accounts [44]. Moreover, widely used dynamic masking interfaces (particularly on mobile browsers [33]) briefly display the most recently typed character before re-masking it, creating routine opportunities for partial or momentary character exposure. Besides UI-level leakage, attackers today can easily acquire identical keyboards, collect training data, and leverage side-channel inference [14], [23]. All this makes mask guessing attacks increasingly practical. Our study therefore provides timely insights into how modern systems behave under partial leakage and how improved modeling of masked passwords can enhance password strength evaluation, inform adaptive throttling and rate-limiting mechanisms, and support the design of more robust user interfaces.

A. Design challenges

In mask guessing, attackers can additionally exploit partial character information from the target password (compared with traditional whole-password guessing), and a given password template (e.g., `**veu4eve*`) usually indicates that the attacker further knows the length and pattern/structure/composition of the target password. This exemplary subtlety partially implies that *the exploration of mask guessing looks deceptively simple, but actually, it is rather challenging*. The following explains why.

Firstly, when applied to mask guessing scenarios, most state-of-the-art password models are unable to assign a probability to the missing character of a template efficiently. As briefly mentioned in [50], [52] and in-depth investigated in this work, character-level generative models like RFGuess [66], FLA [46] and Markov [43], [48] imply a causal order between password characters. This assumption asserts that causality flows in a single and specific direction during the generation process, that is, from the beginning to the end of the string. As a result, these password models are good at assigning probabilities to the following characters by the preceding characters but *not vice versa*. While training companion models with two directions may seem reasonable, it introduces two independent probabilistic models: the forward and reverse-trained models. This approach necessitates careful consideration of additional probabilistic smoothing, which can be complex and non-trivial.

PCFG-based password models [60], [67] are designed for token-based probabilistic modeling. They can only instantiate the basic structures (L/D/S) of passwords with substrings and

cannot infer the missing positions of a password character by character, so these models cannot be directly employed for mask guessing. Deep generative models like PassGAN [29] and CPG [52] have inherent limitations. More specifically, they cannot assign a probability to the generated password, preventing them from sorting the guesses in an optimal order. Thus, a new technical route for mask guessing is needed, but how to design password models to overcome the identified weaknesses has not been systematically explored.

Secondly, it is inherently impractical to directly apply the existing password strength evaluation method, i.e., the Monte Carlo (MC) method [18], to estimate a password’s guess number for mask guessing due to the accuracy vs. efficiency trade-off. Since the way of enumerating large-scale guesses is computationally intensive, previous works [19], [40], [46], [70] have commonly employed the MC method to estimate the guess number of a password under a given probabilistic password model. Its basic idea is to sample from the given password model, i.e., generating random passwords according to the probabilities assigned by the model, and a more accurate result can simply be obtained by increasing the sample size.

However, when applied to mask guessing, the MC method [18] does not provide a guarantee that the randomly sampled passwords will necessarily match a given template (e.g., `**veu4eve*`). It needs to filter out a large number of samples that do not conform to the password template. This may result in an insufficient sample size, which does not guarantee the reliability of the evaluation results.

B. Our contributions

In summary, our contributions are four-fold.

- **Two mask password models.** We propose a Seq2Seq-based password model, namely PassSeq, which can accurately capture the impact of leaked characters on the overall security of the whole password and is applicable to various mask guessing scenarios. *For the first time*, we introduce the Kneser-Ney smoothing technique [38] into password guessing, and propose the Kneser-Ney password model, well mitigating the long-standing issues of overfitting and data sparseness.
- **New guess number evaluation method.** We adjust the generation mechanism of existing character-level autoregressive password models (e.g., Markov [43] and Backoff [43]) from traditional whole-password guessing to mask guessing, and propose a new guess number estimation method based on the Maximum Likelihood Estimation (MLE) technique for mask password guessing. With MLE, we can accurately and efficiently evaluate the password strength under each given template.
- **Extensive evaluation.** By conducting an extensive review of existing side-channel attacks on passwords, we model four different types of realistic mask attackers. The results demonstrate the severe dangers posed by various mask attackers. Particularly, when a targeted attacker knows just *one character* of the victim’s password, the guessing success rate generally doubles within 1-20 guesses (see

Figs. 6(a)-6(d)). Furthermore, within 10 guesses, our PassSeq can boost the effectiveness of existing keystroke acoustic-based side-channel attacks, achieving up to a 167% increase in guessing success rates (see Table IX).

- **Some insights.** We provide a practical application and valuable insights. Specifically, our PassSeq can serve as a reliable password strength meter. Interestingly, mask attackers exploiting PII can achieve a 6%-14% higher guessing success rate when the first character of the target password is exposed, compared to when a randomly chosen character is exposed (see Figs. 6(a)-6(d)).

II. RELATED WORK AND BACKGROUND

In this section, we first introduce major password guessing models and then present our *threat model*.

A. Password guessing attacks

This work divides password guessers in existing research into three types according to different technical routes.

Type-1 research mainly employs *empirical knowledge and heuristic insights* to design various transformation rules (namely mangling rules), such as insertion, reversal, capitalization, and leet (e.g., `password`→`passw0rd`). In 1979, Morris and Thompson [47] pioneered heuristic transformation rules to generate variants of words in the dictionary, and performed password guessing. Besides, there exist some open source password guessing tools (e.g., John the Ripper [5] and Hashcat [25]), taking the mangling rule techniques one step further by leveraging GPUs to automate the guesses at scale. While these heuristics are reasonably successful in practice, they are primarily applied in an ad hoc manner, rather than being constructed from a principled approach.

Type-2 research gradually moves away from the heuristic stage of relying on whimsical ideas and enters the scientific stage based on reliable *probabilistic statistical models*. In 2005, Narayanan and Shmatikov [48] proposed a password model based on the Markov chain, which estimates password probability from every two adjacent characters. In 2014, Ma et al. [43] introduced a number of normalization (e.g., End-symbol) and smoothing (e.g., Backoff) techniques on the Markov chain to overcome the data sparseness and overfitting problems, enabling password generation without being restricted by fixed-length substrings.

In 2009, Weir et al. [67] designed a password guessing model based on probabilistic context-free grammars (PCFG). PCFG [67] views a password as the concatenation of several independent strings of different lengths and character types (i.e., letters, digits, and symbols), corresponding to a specific password structure. Thus, it can generate password structure templates with grammatical rules obtained by statistics, and then fill them with character strings in the training dictionary. Since then, a series of improved PCFG-based methods have been proposed one after another. For example, in 2014, Veras et al. [60] discovered semantic strings in passwords through dictionary matching and classified them into a special category, strengthening the ability to analyze password semantics;

TABLE I
COMPARISON OF MAJOR MASK GUESSING LITERATURE (GAN/WAE=GENERATIVE ADVERSARIAL NETWORK/WASSERSTEIN AUTOENCODER).

Model / Paper	Password length	Probabilistic per-position	Character composition	PII/Reuse-based	Core modeling idea	Notes / Distinction
CPG [52] (Pasquini et al., S&P'21)	Required	Fixed $p=0.5$	✗	✗	Representation learning via GAN/WAE latent smoothness	These models primarily extend trawling password generation to conditional/fixed template-based settings, but lack explicit probabilistic modeling for masked positions and considerations for unknown PW lengths. While Xu et al. [71] and Yang-Wang [72] consider PII/reuse-based scenarios, they are both designed for <i>whole</i> -password guessing rather than mask guessing.
PassBERT [71] (Xu et al., Security'23)	Required	Fixed $p=0.5$	✗	✗	Bi-directional Transformer (pre-train and fine-tune)	
RankGuess-MASK [72] (Yang and Wang, S&P'25)	Required	Fixed $p=0.5$	✗	✗	Adversarial ranking with offline reinforcement learning	
PassSeq and Kneser-Ney (This work)	Unknown L considered	Fixed $p=0.5$; Side channel-aware	✓	✓	Seq2Seq with attention and Kneser-Ney smoothing	We quantify four mask guessing scenarios that incorporate side-channel priors, PII, and previously leaked (sister) passwords.

In 2015, Houshmand et al. [32] added the keyboard and multiword patterns to PCFG to enhance its ability to crack weak passwords. In 2022, Wang et al. [61] investigated an identification method for password segments and introduced the ReSeg-PCFG model. Their study revealed that the number of segments significantly influences password security. In 2024, Huang et al. [34] integrated two prominent password guessing models, PCFG [67] and the Markov-based OMEN [20], with Hashcat to fully leverage GPU acceleration, significantly improving the guessing efficiency of both models.

Type-3 research is mainly based on *machine/deep learning*, taking more advantage of large-scale data to address the data sparseness and overfitting problems existing in traditional statistical password models. In 2017, Melicher et al. [46] introduced long short-term memory networks [30] to password guessing and proposed FLA (Fast, lean, and accurate). Like Markov [43], FLA is also trained to generate the next character of a password given the preceding characters. The difference is that FLA is able to automatically assign a small (but non-zero) probability to each character in the alphabet, instead of using smoothing techniques. Recently, Pasquini et al. [51] proposed a FLA-based universal password model, which relies on auxiliary data (e.g., email) to build a pre-trained model for the target group of users. In addition, there are several leading targeted password models, such as TarGuess [64], which can exploit victims' PIIs, and Pass2Path [49], which focuses on credential stuffing/tweaking attacks.

In 2021, Pasquini et al. [52] successfully mitigated the training collapse problem of PassGAN [29]. They achieved this by implementing a form of stochastic smoothing over the representation of strings, leading to a substantial improvement in the performance of GAN-based approaches. On this basis, they constructed two guessing frameworks, i.e., conditional password guessing (CPG) and dynamic password guessing (DPG). They also showed that existing password models that are designed for traditional whole-password guessing (e.g., PCFG [67], Markov [48], and FLA [46]) can be used for mask guessing by filtering the generated guesses through templates, but their guessing success rates are not high. Moreover, this approach would generate a lot of redundant guesses that do not conform to the template, which might be inefficient.

In 2023, Xu et al. [71] introduced pre-training and fine-tuning techniques for password guessing, and proposed PassBERT, a framework based on Transformer. The authors first pre-trained a password model with the masked language modeling objective using large-scale datasets. Then, they designed attack-specific fine-tuning approaches to tailor the pre-trained password model to three attack scenarios.

In 2025, Yang and Wang [72] proposed RankGuess, a password guessing framework based on adversarial ranking. By formulating the guessing task as a Markov Decision Process, RankGuess demonstrates improved performance over existing models across three major guessing scenarios (Trawling, PII-based, and Mask guessing scenarios). Recently, several large language model-based password models (e.g., PassGPT [53], PagPassGPT [55], and PassLLM [76]) have been proposed. However, these models are focused on traditional whole-password guessing, rather than mask guessing attackers that can exploit the victim's partial password information (e.g., length and/or some characters) through side-channel cues.

Summary. To the best of our knowledge, only a few studies have explicitly addressed the mask password guessing scenario, including the conditional password guessing (CPG) framework [52], PassBERT [71], and RankGuess-MASK [72]. However, these works focus mainly on canonical settings, where part of the password is already known or structurally fixed. *They overlook more realistic and practical variants*, such as side-channel attacks (see Fig. 10) that provide probabilistic information about character-level distributions (which we summarize and compare in Table I). Fundamentally, mask password guessing can be viewed as a problem of probability allocation/smoothing over missing characters. This raises an important question: Could traditional statistical smoothing techniques (e.g., Backoff [43]) be more suitable than complex neural architectures in this context? Furthermore, are there more effective smoothing techniques that better align with the nature of mask guessing? These questions remain largely unexplored in existing literature.

B. Threat model

We introduce a formal attacker model for real-world mask guessing scenarios, grounded in the partial information that can be exploited by the attacker. Based on the characteristics

TABLE II
SUMMARY OF REPRESENTATIVE SIDE-CHANNEL CHARACTERISTICS FOR CONTEXTUALIZING OUR EVALUATED MASK-GUESSING SCENARIOS.

Attack types and base success rates / prevalence / practical realism	Key experimental settings / parameters of representative studies	Experimental setup in this work
Shoulder-surfing [7], [45] ① Using video recordings of participants unlocking devices, Aviv et al. [7] report shoulder-surfing success rates of 34.9%/10.8% and 56.7%/26.5% for 4-digit/6-digit PINs after one and multiple observations. ② An anonymity-preserving survey by Marques et al. [45] finds that about 31% of participants admit to having looked through someone else's phone without permission within the past year. When weighted to the U.S. adult population, this corresponds to roughly 20%, with the behavior most prevalent among younger adults (18-24 years, $\approx 52\%$) and smartphone owners. ③ Dynamic masking [†] in major Android/iOS browsers makes partial character exposure a realistic threat and directly motivates our experimental scenarios [33].	① Aviv et al. [7] simulated shoulder-surfing attacks by showing participants pre-recorded videos of users entering 4- and 6-digit PINs from different viewing angles ($N = 1,173$). ② Using anonymous online list experiments ($N \approx 2,000$) via Google Consumer Surveys, Marques et al. [45] measured how often people looked through others' phones without permission. ③ Hu et al. [33] investigated password masking practices across popular authentication systems, including the Google CrUX Top-1K websites and major browsers (e.g., Safari and Samsung).	We consider nine partial password-leakage cases covering scenarios where the attacker exploits: (a) personally identifiable information (PII) of the target, (b) previously leaked passwords of the same user, and (c) no related information. More specifically, these include: (1) PII-based scenarios with length-only leakage; (2) PII-based scenarios with leakage of length and one character; (3) Reuse-based scenarios with length-only leakage; (4) Reuse-based scenarios with leakage of length and one character; (5) scenarios with leakage of the first character without length; (6) scenarios with leakage of length and the first character; (7) scenarios with leakage of 50% characters without length; (8) scenarios with leakage of 50% characters with length; (9) scenarios with length-only leakage. Results for the PII/reuse-based scenarios appear in Figs. 6 and 13, and results for the trawling scenarios are in Table VII and Fig. 7.
Bicycle attack (length-leak) [24] About 84% of Alexa Top-100 websites are vulnerable to password-length leakage [24].	Password lengths are inferred from TLS request sizes of Alexa Top-100 websites (e.g., Gmail), revealing the exact length through the correlation between plaintext and ciphertext sizes [24].	The risks introduced by password length leakage are quantitatively evaluated in our constructed scenarios (1), (3), and (9) above; see Figs. 6 and 13 for PII and Reuse-based mask guessing scenarios and Figs. 7 and 15 for trawling guessing scenarios.
Keystroke audio feedback [14], [23] ① Under a threat model where an attacker records keystrokes using a nearby smartphone or through a remote voice call (e.g., Zoom or Skype), Harrison et al. [23] extract mel-spectrogram features and train a neural classifier that achieves up to 95% <i>character-level</i> accuracy for smartphone-recorded keystrokes. ② Cardaioli et al. [14] show that inter-keystroke timings from keypad audio markedly strengthen PIN-guessing attacks, boosting five-attempt success rates from $\approx 44\%$ (with only the pressed-digit set) to $\approx 89\%$ when combined with typing behavior, one known digit, and the pressed-digit set.	① Data were collected on a MacBook Pro 16-inch (2021) using 36 keys (0-9, a-z; each pressed 25 times). Recordings were captured by an iPhone 13 mini placed 17 cm from the keyboard in stereo at 44.1 kHz and 32-bit [23]. ② The dataset contains audio recordings of 4-digit PIN entries from 22 users on a simulated ATM keypad in a noisy indoor environment (SNR -15 dB), captured at 48 kHz from a distance of about 1.5 m. Each keystroke produced uniform audio feedback with millisecond-level timestamps.	We follow the threat model of Harrison et al. [23] and use per-character probability distributions inferred from keystroke audio as priors for our PassSeq. Note that Harrison et al. [23] collected repeated presses of the same physical key rather than actual password sequences. To construct a more realistic guessing scenario, we collected over 1,500 keystroke samples of real PIN entries from 16 users across 11 popular keyboards (see Table IX). Besides, we employ the method of [3] to enhance the authors' open-source pipeline, improving its character-level accuracy on our collected dataset from approximately 40% to 94%. Cardaioli et al. [14] combine inter-keystroke timing information with knowledge of one PIN digit and/or the unordered set of digits composing the PIN. Following their threat model, we integrate our PassSeq password model with the same forms of leaked information. (see Fig. 5).
Thermal residue [4], [6] ① Alotaibi et al. [6] show that, using AI-driven analysis of thermal images taken after password entry, an attacker can achieve success rates ranging from 92% (6-symbol) to 55% (16-symbol). When thermal images are captured shortly after entry, attackers achieve 86% at 20s, 76% at 30s, and 62% at 60s. ② Abdelrahman et al. [4] demonstrate that thermal residue on smartphone screens can be exploited to reliably recover recently entered PINs. Their Thermal-Analyzer achieves a success rate of 72–100% when attacks were performed within 30s of authentication, with accuracy dropping sharply when the thermal traces were observed 45-60s after entry.	① Thermal images were collected from 21 participants who entered passwords of varying lengths (6, 8, and 12-16 characters) on a Microsoft Wired Keyboard 600 with ABS keycaps. An Optris PI 450 camera was used to capture the thermal snapshots at 20/30/60s [6]. ② 18 participants entered multiple 4-digit PINs on Samsung Galaxy Note Edge smartphones. Thermal images were captured using an Optris PI450 thermal camera positioned ≈ 80 cm from the phone, at 0-60s after entry [4].	We construct per-position disclosure probabilities inspired by the position-dependent recovery probabilities observed in ThermoSecure's decoding process (see Fig. 11 and Table I of [6]). Specifically, we define a simple six-position recognizability profile, $[0.17, 0.20, 0.25, 0.33, 0.50, 1.00]$, which captures the monotonic position-wise increase in decoding probabilities observed in thermal-residue analysis. These values serve as the base disclosure rates in our probabilistic leakage guessing scenarios (see Fig. 4). To extend this profile to arbitrary password lengths L , we map the six reference positions onto the normalized interval $[0, 1]$ and apply linear interpolation at L equally spaced indices, yielding a smooth, monotone recognizability curve. To account for moderate variability (approximately 5-10%) arising from temperature decay and imaging delays, we add Gaussian perturbations $\mathcal{N}(0, 0.05^2)$ to each interpolated value q_i and clip the results to $[0, 1]$.
Smudge attacks [16] Within 20 attempts, the smudge-supported guessing attack achieves $\approx 74\%$ success with clean residue and $\approx 32\%$ with noisy/obscured residue [16].	Automated extraction of smudge traces from smartphone photos (using Samsung Galaxy S4 with train/test = 219/93, $N = 12$); evaluated unlock, call, and social-app scenarios [16].	Smudge attacks (an example is shown in Fig. 1) and thermal residue (e.g., [4]) reveal the set of characters composing the password but not their exact positions. We model this scenario using PIN data in Fig. 5 and quantify the incremental risk introduced by such partial composition leakage.

[†] All major mobile browsers on Android and iOS implement dynamic masking (where a recently typed character is unmasked, but previously-typed characters are masked; see Fig. 1 in [33]), allowing shoulder-surfing attackers to observe certain visible characters during password entry.

and granularity of the available information, we identify four generalized categories of attackers. Each attacker is characterized by information K , which constrains an effective guessing space $\mathcal{G}_K$ and informs the guessing strategy π_K .

General setting. Let $\mathcal{P}$ denote the password space and $\text{pw}^* \in \mathcal{P}$ the ground-truth password of length L . The attacker attempts to guess pw^* given partial prior knowledge K . We define the constrained candidate space as:

$$\mathcal{G}_X = \{\text{pw} \in \mathcal{P} \mid \text{pw satisfies } K\}, \quad (1)$$

and the goal is to construct a guessing policy π_X that maxi-

mizes the success rate under a given guessing budget.

Type I: Template attackers ($\mathcal{A}_{\text{template}}$). These attackers possess an exact mask template PT, e.g., `**veu4eve*`, where known positions are revealed and the others remain masked. The corresponding candidate set is:

$$\mathcal{G}_{\text{template}} = \{\text{pw} \in \mathcal{P} \mid \text{pw}[i] = \text{PT}[i] \forall i \text{ s.t. } \text{PT}[i] \neq *, \quad (2)$$

where c_i is the character at position i in the password, and $*$ is the masked character. Such templates can be obtained from shoulder-surfing (e.g., [22], [39]) and bicycle attacks [24].

TABLE III
ATTACKER CAPABILITIES CONSIDERED IN THIS WORK.

Mask attacker types	Password length	Character information with 100% certainty	Character information with some probability [†]	Character compositions [‡]	PII*	Leaked passwords*	Existing literature	Experimental results
$\mathcal{A}_{\text{template}}$	✓	✓					[52], [71], [72]	Table VII
$\mathcal{A}_{\text{stat}}$	✓		✓	✓			This work	Fig. 4, Fig. 5, Table IX
$\mathcal{A}_{\text{context}}$	✓	✓			✓	✓	This work	Fig. 6, Fig. 13
$\mathcal{A}_{\text{length}}$		✓					This work	Fig. 6, Fig. 7

[†] Keystroke audio feedback (e.g., [9], [10], [14]) or thermal residue (e.g., [4], [6]) can infer the likelihood of password characters (with position determined).

[‡] Smudge attacks (e.g., [8] and Fig. 1) are a type of side-channel attack that can infer the composition of password characters (with position undetermined).

*PII=Personally Identifiable Information. Besides the partial information of the victim's password (e.g., password length and/or some characters), the mask attacker can exploit the victim's name, birthday, and previously leaked passwords to further enhance their guessing success rates.

In response to this attack scenario, Pasquini et al. [52], Xu et al. [71], and Yang and Wang [72] have respectively proposed the CPG, PassBERT, and RankGuess password models. These works employ various password templates to generate large-scale password guesses that conform to the templates, and then evaluate their password models on specific (template-conformed) test sets (that are different from the training sets). While all three studies highlight that some real-world side-channel attacks could obtain partial password information, they do not clearly articulate which specific parts of a password can be obtained under which types of attack scenarios.

To comprehensively evaluate the practical risks of the mask scenario, we extensively review existing studies that can obtain partial password information and further classify the mask scenarios to better align with real-world situations. Still, this basic mask scenario (i.e., the $\mathcal{A}_{\text{template}}$ attacker) will serve as a benchmark for comparing the advancements of the models proposed in this paper (see Table VII).

Type II: Distributional or compositional attackers ($\mathcal{A}_{\text{stat}}$). These attackers acquire statistical or compositional side-channel information about the target victim's password. One form of leakage is a positional character distribution $D = \{D_i\}_{i=1}^L$, where each D_i denotes the distribution over possible characters at position i . Specifically,

$$D_i(c) = \Pr(\text{pw}^*[i] = c), \quad (3)$$

which denotes the probability that character c appears at position i of the target password pw^* . Such distributions are typically derived from thermal residue on keyboards (e.g., [4], [6]) or keystroke audio feedback (e.g., [9], [14], [74]). Another common form is an unordered character multiset $M \subseteq \Sigma^*$, as revealed by smudge attacks on touchscreens (see Fig. 1), which disclose the set of characters used in the password without any positional information. In both cases, the attacker's guessing strategy involves either maximizing the joint likelihood of characters under D or enumerating permutations of M that form valid password candidates.

Type III: Contextual prior attackers ($\mathcal{A}_{\text{context}}$). These attackers enhance traditional mask guessing by incorporating user-specific prior knowledge into the attack process. Such contextual priors include personally identifiable information (PII), denoted as $\mathcal{I}$ (e.g., the victim's name, birthday, or email address), as well as previously leaked passwords, denoted as $\mathcal{R}$. By conditioning guesses on $\mathcal{I}$ or adapting patterns from $\mathcal{R}$, these attackers can significantly narrow the search space and increase the guessing success rate.

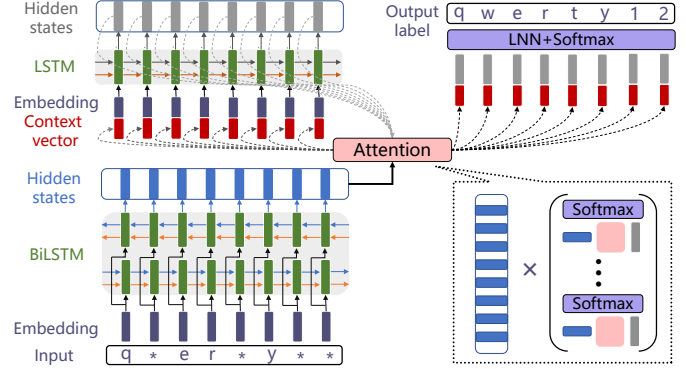


Fig. 2. The network architecture of our Seq2Seq. Here we show a high-level example of the architecture with one layer of LSTM/BiLSTM.

Type IV: Uncertain-length template attackers ($\mathcal{A}_{\text{length}}$). These attackers obtain partial template information PT but cannot determine the exact password length. Instead, a length candidate set $\mathcal{L} \subseteq \mathbb{N}$ is given. The attacker must consider multiple candidate lengths in parallel:

$$\mathcal{G}_{\text{length}} = \bigcup_{\ell \in \mathcal{L}} \{\text{pw} \in \mathcal{P} \mid |\text{pw}| = \ell \text{ and pw satisfies PT}\}. \quad (4)$$

Table II provides representative side-channel attacks (e.g., shoulder surfing [7], [45], keystroke audio feedback [14], [23], and thermal residue [4], [6]) and their key characteristics (including baseline success rates, prevalence, practical feasibility, and typical experimental settings) which contextualize our evaluated mask-guessing scenarios. The corresponding attacker capabilities are listed separately in Table III.

III. OUR TWO NEW SEQUENCE PASSWORD MODELS: PASSSEQ AND KNESER-NEY

We first elaborate on our PassSeq model, then introduce the Kneser-Ney smoothing technique [38] to password guessing, and finally present a new method for estimating the guess number required for mask guessing.

A. Our PassSeq password model

Mask guessing scenarios can be modeled as *character-level sequence-to-sequence language modeling tasks* that leverage prior information, such as PII (if available) and template sequences, to predict target password sequences. The Seq2Seq model, with its encoder-decoder architecture [56], offers significant structural flexibility, making it well-suited for this task. More specifically, the encoder can be composed of bidirectional LSTM [30], capturing bidirectional information

of the password templates. This is particularly useful for templates where the continuity of local n -grams is disrupted (e.g., $*2*45*as*$). As we qualitatively analyze in Section IV-B, bidirectional encoding more effectively accommodates such templates than traditional n -gram smoothing. Further, the decoder structure ensures that passwords are generated character by character in an *autoregressive* manner from left to right (see equation 5), aligning better with the typical human behavior when creating and entering passwords. This also results in the model outperforming the PassBERT model [71] in mask scenarios. While transformer architectures [59] also perform well in bidirectional sequence tasks, the simplicity of the Seq2Seq model aligns with Occam’s razor principle, making it a pragmatic choice for our application.

Our Seq2Seq network architecture is shown in Fig. 2. The encoder converts the input password character tokens to a representation vector and then feeds it to the decoder to generate output tokens based on a target vocabulary. The initial Seq2Seq model was proposed by Sutskever et al. [56], and originally applied to the English-French translation task. In 2015, Luong et al. [42] employed an attention mechanism, which makes use of *all* the encoder outputs and thereby improves the model performance. Thus, we also employ the attention mechanism to build our neural network.

When applying Seq2Seq for mask guessing scenarios, the password template PT (e.g., $q*er*y**$) is encoded as the input of the model, and the corresponding password pw (i.e., $qwerty12$) is used as the output label. We define the probability of predicting the complete password $pw(=c_0, \dots, c_n)$ given the corresponding template PT as

$$\Pr(pw|PT) = \prod_{i=1}^n \Pr(c_i|c_0, \dots, c_{i-1}, PT). \quad (5)$$

In the training phase, we use cross-entropy to quantify the discrepancy between the outputs and the corresponding labels. In the generation phase, we can use beam search or our proposed mask search algorithm (see Algo. 1) to generate guesses. We call the resulting model PassSeq, and put the details of the model structure in Table XIX. We note that Pal et al. [49] and Xiu and Wang [69] have respectively proposed leading password reuse models based on Seq2Seq. However, their model architectures, training/generation paradigms, and core application scenarios are all different from our PassSeq.

B. Our Kneser-Ney password model

Traditional statistical password models (e.g., Markov [48]) generally rely on smoothing techniques to address the data sparsity issue inherent in n -gram modeling. In 2014, Ma et al. [43] incorporated the backoff smoothing method [36] into the Markov password model [48], resulting in the widely used Backoff password model. While effective in some scenarios, backoff smoothing tends to assign overly simplistic probabilities in the absence of high-order n -gram matches, limiting its generalization ability in sparse or unseen contexts. To better address these limitations, we explore the use of

Kneser-Ney (KN) smoothing, a technique that has shown remarkable robustness in natural language modeling but remains underexplored in password modeling. To motivate this choice and highlight the shortcomings of traditional approaches, we provide a detailed toy example illustrating the differences between MLE, backoff smoothing, and Kneser-Ney smoothing in character-level prediction in the Appendix A.

Model formulation. In practice, KN smoothing has emerged as one of the most theoretically grounded and empirically robust smoothing methods in language modeling, owing to its ability to balance count discounting and context diversity [27], [38], [54]. Given the analogous challenges in character-level password modeling (e.g., extreme sparsity and high contextual variability), we incorporate KN smoothing as a principled and effective solution within our probabilistic framework. Let the password substring be $pw_i^j = c_i c_{i+1} \dots c_j$. The conditional probability of the next character c_i given context pw_{i-n+1}^{i-1} is:

$$\Pr_{KN}(c_i | pw_{i-n+1}^{i-1}) = \frac{C(pw_{i-n+1}^{i-1}) - D(C(pw_{i-n+1}^{i-1}))}{\sum_{c_i} C(pw_{i-n+1}^{i-1} c_i)} + \gamma(pw_{i-n+1}^{i-1}) \cdot \Pr_{KN}(c_i | pw_{i-n+2}^{i-1}) \quad (6)$$

The discount function $D(x)$ is:

$$D(x) = \begin{cases} 0, & \text{if } x = 0 \\ x - \frac{n_1}{n_1 + 2n_2} \cdot \frac{n_x + 1}{n_x} (x + 1), & \text{if } x = 1, 2, 3 \\ D(3), & \text{if } x > 3 \end{cases} \quad (7)$$

where n_x is the number of n -gram substrings with frequency x . The backoff weight $\gamma(\cdot)$ is:

$$\gamma(pw_{i-n+1}^{i-1}) = \frac{D(1)N_1(pw_{i-n+1}^{i-1}) + D(2)N_2(pw_{i-n+1}^{i-1})}{\sum_{c_i} C(pw_{i-n+1}^{i-1} c_i)} + \frac{D(3)N_{3+}(pw_{i-n+1}^{i-1})}{\sum_{c_i} C(pw_{i-n+1}^{i-1} c_i)} \quad (8)$$

Here, N_1 , N_2 , and N_{3+} denote the number of distinct continuations of pw_{i-n+1}^{i-1} appearing exactly once, twice, and three or more times, respectively.

C. Our guess number estimation method

Considering that enumerating large-scale guesses is computationally intensive, previous work (e.g., [40], [46], [70]) generally employ the Monte Carlo (MC) method [18] to estimate a password’s guess number (given a password probabilistic model) if the guesses are in descending order of likelihood. However, MC is not suitable for mask guessing, since it lacks an efficient way to ensure that the sampled passwords match a given template. To address this limitation, we consider modifying the original MC estimation as follows:

$$C_\Delta = \sum_{pw \in \Theta} \begin{cases} \frac{1}{n \cdot p(pw)} & \text{if } p(pw) > p(pw^*), \\ 0 & \text{otherwise.} \end{cases} \quad (9)$$

Here, Θ is a sample of size N from the password space Γ , and C_Δ estimates the number of passwords more likely than pw^* . By constraining pw to also match a given password template PT^* , we can refine the estimation to match mask guessing scenarios. However, this adaptation is often impractical. Even with 10^7 samples from a trained Markov model

[43], the average number of passwords matching a specific template remains low (e.g., 0-16), far below the recommended $N = 10^4$, leading to unreliable estimation.

Instead, we propose using Maximum Likelihood Estimation (MLE) to estimate password strength in mask guessing. Given that the number of passwords matching a template PT^* is finite and known (denoted n), and that m of them satisfy $p(pw) > p(pw^*)$, we aim to estimate m .

To this end, we perform t independent sampling trials. In each trial, we uniformly sample N passwords from the set of candidates matching PT^* , and count how many of them satisfy $p(pw) > p(pw^*)$. Let x_i denote the number of such passwords in the i -th trial. The MLE of p is then computed as:

$$\tilde{p} = \frac{1}{tN} \sum_{i=1}^t x_i. \quad (10)$$

This estimator is unbiased, and its variance satisfies $\text{Var}(\tilde{p}) = \mathcal{O}(1/(tN))$, indicating rapid convergence as the number of trials and sample size increase.

Full derivation and theoretical analysis of MLE (including variance and convergence rate) can be found in Appendix D.

Intuitively, MLE follows a principle similar to the *mark-recapture* estimation: by uniformly sampling n passwords from N total passwords (within a given template) and observing m passwords that are stronger than the target password, we can estimate its rank as $M = mN/n$. Repeating this process multiple times (e.g., 4-10) can produce an (approximately) unbiased estimate of the target password’s rank, thereby aligning with the attacker’s ranking strategy.

In practice, MLE is particularly suitable when estimating large-scale guess numbers for a given password under a given template (e.g., $\geq 10^8$ guesses). We provide a brief analysis of the computational overhead introduced by MLE in Section IV-F, highlighting its significant efficiency advantage over directly generating and ranking the entire candidate space to determine the guess number. The large-scale experimental results reported in Table X (in Appendix B) are all obtained using our MLE-based estimation.

IV. EXPERIMENTS

Now we first elaborate on our password datasets and then fairly/comprehensively evaluate our proposed PassSeq and Kneser-Ney with their foremost counterparts (i.e., Markov [43], Backoff [43], Hashcat [25], PCFG [67], OMEN [20], CPG [52], PassBERT [71], and RankGuess [72]) in typical mask guessing scenarios.

A. Our datasets

Datasets. We introduce 15 real-world datasets with over 3.4 billion passwords (see Table IV), including eight from English sites, six from Chinese sites, and one mixed (COMB, containing passwords from over 50 countries/regions). Among them, three datasets (i.e., 12306, Rootkit, and Clixsense) are associated with various PII (e.g., name, email, and birthday). For PII or reuse-based guessing scenarios, we obtain

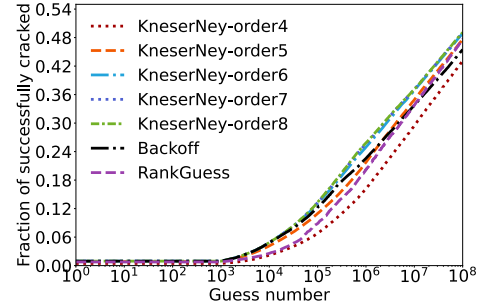


Fig. 3. We evaluate our Kneser-Ney password model with different n -gram orders ($n=4, 5, 6, 7, 8$), alongside Backoff [43] and RankGuess [72] baselines. Results show that the guessing success rate of Kneser-Ney is slightly higher when $n=7/8$. Therefore, we set $n=7$ for all subsequent experiments.

another two PII datasets (see Table V) and four password-pair datasets by matching email (see Table VI). These datasets are compromised by hackers or leaked by insiders, and have been publicly available for some time. Following the data cleaning method employed by [43], [63], [72], we remove non-password strings in the original dataset, including headers, descriptions, footnotes, and non-ASCII strings. We also remove passwords with length > 30 , because they are likely to be constructed by password managers or just junk information.

B. Template-based mask guessing scenarios

We compare our PassSeq and Kneser-Ney with eight state-of-the-art password guessers in typical mask guessing scenarios (i.e., the type-I attacker $\mathcal{A}_{\text{template}}$), and they are PCFG [67], OMEN [20], Markov [43], Backoff [43], Hashcat [25], CPG [52], PassBERT [71], and RankGuess [72]. The reasons for selecting these password models are as follows: PassBERT [71], CPG [52] and RankGuess [72] are state-of-the-art password models *specifically designed for* typical mask guessing scenarios; PCFG [67] remains the most influential structure-based guesser; Hashcat’s mask mode is a widely used attack method among real-world attackers; Markov/OMEN/Backoff are leading character-level autoregressive models that rely on statistical smoothing techniques. Particularly, we adapt Markov and Backoff to mask guessing scenarios using our proposed generation algorithm (see Algorithm 1).

We make sure the eight approaches work on the same training (i.e., 80% Rockyou as with [52]) and test sets, and manage to use/obtain their codes shared/open-sourced by the original authors. For all model parameters, we follow the best recommendations. The specific setup is as follows.

- **PassSeq.** For the built-in Seq2Seq model, we set the context length to 16, the hidden size, and the embedding size to 256 and 128, respectively. See Table XIX (in Appendix) for detailed parameters/model structure.
- **PCFG.** We use the latest open-source implementation of the PCFG cracker [67] and follow the default grammar extraction and generation pipeline.
- **OMEN.** We use the public release of OMEN [20] and execute it with its recommended default configuration.
- **Kneser-Ney.** We set the order to 7 (i.e., 8-gram) based on our experimental results (see Fig. 3).

TABLE IV
BASIC INFORMATION ABOUT OUR 12 PASSWORD DATASETS (PW=PASSWORD).

Dataset	Language	Service type	Leaked time	Original size	Non-PWs/Non-ASCII	PW length>30	Removed	After cleaning	Remaining
Post Millennial	English	News outlet	May 2024	38,902	0	31	0.08%	38,871	99.92%
Wishbone	English	Mobile app	Jan. 2020	10,092,037	798	250	0.01%	10,090,989	99.99%
000Webhost	English	Web hosting	Oct. 2015	15,299,907	69,606	1,131,721	7.76%	14,098,263	92.24%
LinkedIn	English	Job hunting	July 2012	54,656,615	0	17,157	0.03%	54,639,458	99.97%
Yahoo	English	Portal	July 2012	453,391	10,709	1	2.37%	442,681	97.63%
Rockyou	English	Social forum	Dec. 2009	32,603,387	149,971	2,068	0.47%	32,451,348	99.53%
Taobao	Chinese	E-commerce	Feb. 2016	15,072,418	0	86	0.00%	15,072,332	100.00%
Sohu	Chinese	Portal	Oct. 2015	14,755,046	147,429	50	1.00%	14,607,567	99.00%
126	Chinese	Email	Dec. 2011	6,392,568	0	599	0.00%	6,391,969	100.00%
CSDN	Chinese	Programmer forum	Dec. 2011	6,426,872	0	125	0.00%	6,428,285	100.00%
Dodonew	Chinese	E-commerce	Dec. 2011	16,282,276	4,975	12,070	0.10%	16,265,231	99.90%
COMB	Mixed	Mixed	Feb. 2021	3,279,064,312	14,948,181	10,576,452	0.78%	3,253,539,679	99.22%

TABLE V
BASIC INFORMATION OF OUR PII DATASETS.

Dataset	Language	Items num	Types of PII useful for this work
12306	Chinese	129,303	Email, User name, Name, Birthday, Phone
Dodonew-PII	Chinese	80,758	Email, User name, Name, Birthday, Phone
000Webhost-PII	English	79,580	Email, User name, Name, Birthday
Rootkit	English	69,418	Email, User name, Name, Birthday
ClixSense	English	2,222,045	Email, User name, Name, Birthday

TABLE VI
BASIC INFORMATION OF OUR PASSWORD REUSE DATASETS.[†]

Language	Training set setup	Size (pairs)	Test set setup	Size (pairs)
Chinese	CSDN→126	97,915	CSDN→12306	12,635
Chinese	CSDN→12306	12,635	CSDN→126	97,915
English	000Web.→Clixsense	150,273	000Web.→Yahoo	36,936
English	000Web.→ClixSense	150,273	000Web.→LinkedIn	231,452

[†] $A \rightarrow B$ means that a user's password at service A can be exploited by an attacker to help attack this user's account at service B . 000Web.=000Webhost.

- **Markov/Backoff.** For Markov, we set the order to 3, and adopt the Laplace and End symbol technique recommended by [43]. For Backoff, we set the count threshold to 10 as recommended by [43].
- **Hashcat.** We employ the mask attack mode with the given password templates (Markov optimization).
- **CPG.** We use the source code of CPG [52] (i.e., PasswordAE) with the default parameters.
- **PassBERT.** PassBERT [71] provides two pre-trained variants: one on natural language and one on password data. We use the latter due to its higher cracking rates.
- **RankGuess-MASK.** We run the authors' provided code with the recommended hyperparameters [72]. The model is trained using KL divergence as the loss function and the Adam optimizer (learning rate=0.001, $\beta_1=0.5$, $\beta_2=0.999$). A penalty factor $\gamma=0.1$ is applied, with 17 groups and a dropout rate of 0.3 for both the guesser and ranker.

At IEEE S&P'21, Pasquini et al. [52] evaluated existing password models (e.g., Markov [43] and FLA [46]) in mask guessing scenarios by filtering the generated guesses through templates. However, they also revealed that this solution has two main drawbacks: (1) It's costly and storage-demanding; (2) It's difficult for such models to generate relatively low-probability guesses. For a fair comparison, we first closely follow the experimental setup in [52], [71], [72] (e.g., the following template classification method is exactly the same as [52], [71], [72]). Then, we propose a generation algorithm (see Algorithm 1) and adapt the existing character-level au-

toressive password models (including Markov [43], [48] and Backoff [43]) for mask guessing.

More specifically, we first construct a test set consisting of password templates, where the character set includes 94 printable ASCII characters (excluding space) and a wildcard/mask symbol *. Each password template (PT) is extracted from a randomly sampled password in a validation set V . We use LinkedIn as the validation set and retain only passwords with length ≤ 30 , resulting in 6×10^6 unique passwords.

For each template, we uniformly sample a password pw from V and independently mask each character in pw with probability $p=0.5$ (the rationale for this choice is discussed later). We then extract all passwords in V that conform to a given PT, forming the candidate set V^{PT} . Depending on the size of V^{PT} , the templates are divided into four classes.

- $T_{\text{common}} = \{PT \mid \forall PT, |V^{PT}| \in [1, 000, 1, 500]\}$
- $T_{\text{uncommon}} = \{PT \mid \forall PT, |V^{PT}| \in [50, 150]\}$
- $T_{\text{rare}} = \{PT \mid \forall PT, |V^{PT}| \in [10, 15]\}$
- $T_{\text{super-rare}} = \{PT \mid \forall PT, |V^{PT}| \in [1, 5]\}$

Each class contains 30 templates. This setup corresponds to the $\mathcal{A}_{\text{template}}$ attacker, that is, the attacker knows the length and some characters (with certainty) of the target password. Essentially, template-based mask guessing inherently relies on the mask probability p to generate the corresponding password templates. Choosing a value of p that is too high or too low may bias the evaluation of different models in mask guessing scenarios. More specifically, a *high* probability (e.g., $p=0.8$, as in `ab*****`) produces too many masked positions and makes the task closely resemble whole-password guessing, while a *low* probability (e.g., $p=0.2$, as in `aaaaaaaa**`) results in an extremely small search space that can be nearly exhausted by simple enumeration (generally within 10^6 guesses, e.g., the total space of a template with three masked characters is only 94^3). Therefore, we adopt $p=0.5$, which is consistent with prior work (e.g., CPG [52], PassBERT [71], and RankGuess-MASK [72]) and provides a balanced and realistic evaluation setting. In addition, we introduce three masking strategies and report their comparative results in Appendix C. Overall, variations in masking strategies have only a marginal effect on guessing success rates, and $p=0.5$ is a simple yet effective setting for our study.

Note that Hashcat [25], CPG [52], PassBERT [71] and RankGuess-MASK [72] can be directly applied to mask

TABLE VII

PROPORTION OF AVERAGE MATCHED PASSWORDS OVER THE BIASED PASSWORDS TEST-SET DIVIDED INTO FOUR CLASSES (RG=RANKGUESS).[†]

Templates class	Markov-d [43]	Markov-m	Backoff-d [43]	Backoff-m	Hashcat [25]	PCFG [67]	OMEN [20]	CPG [52]	RG-MASK [72]	PassBERT [71]	KN-m	PassSeq
Common [1,000-1,500]	0.2806	0.6457	0.4753	0.8179	0.2803	0.6104	0.3337	0.6312	0.6551	0.4548	0.8366	0.7760
Uncommon [50-150]	0.1041	0.4624	0.4720	0.8038	0.1832	0.6291	0.2123	0.6539	0.5616	0.3538	0.8823	0.7957
Rare [10-15]	0.1709	0.6162	0.5273	0.8879	0.1823	0.8441	0.2590	0.6601	0.7193	0.4633	0.9244	0.8947
Super-Rare [1-5]	0.0580	0.2830	0.1656	0.4094	0.1066	0.1902	0.0285	0.2783	0.3147	0.1931	0.4270	0.4879

[†] The evaluation metric is the average guessing success rate across 30 password templates within each class. The guessing success rate for each template is calculated as $N_{\text{intersection}} / |V^{\text{PT}}|$, where $N_{\text{intersection}}$ represents the intersection between the candidate passwords and the passwords in the validation set (i.e., V^{PT}). “-d” means default mode: We directly generate 10^9 candidate passwords using the original/default generation methods of these models (These generated passwords generally include a significant number that does not match the provided template), and then calculate the guessing success rate; “-m” means mask mode: We use Algorithm 1 to generate 10^6 candidate passwords that naturally match each provided template.

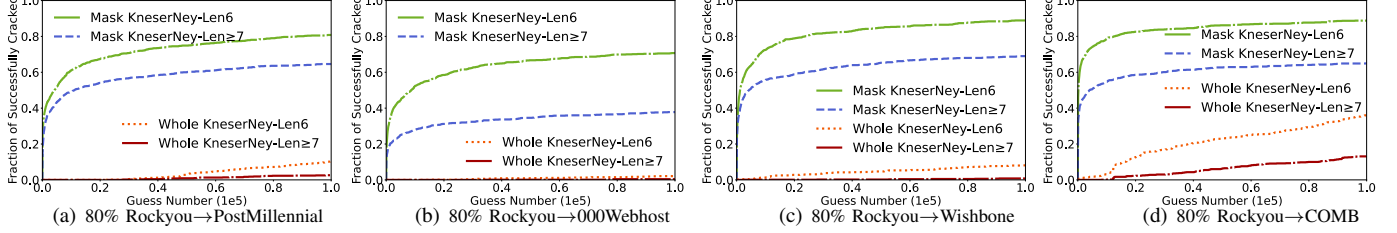


Fig. 4. Mask password guessing under probabilistic character leakage (i.e., the type-2 attacker $\mathcal{A}_{\text{stat}}$). We adopt a thermal-residue-based leakage model in which per-position leakage follows a front-to-back increasing key probability profile, inspired by Alotaibi et al. [6] (Fig. 11). The positional leakage probabilities are set to $[0.17, 0.20, 0.25, 0.33, 0.50, 1.00]$ for length-6 passwords and are extended to other lengths via linear interpolation. For each dataset, we randomly sample 1,000 passwords and construct corresponding probabilistic templates. Each point (X, Y) in the figure represents that, at X guesses, a proportion Y of templates have their original passwords successfully recovered. Under such probabilistic leakage, the guessing success rate of $\mathcal{A}_{\text{stat}}$ is approximately 2.47-33.67 times (within 10^5 guesses) higher than in scenarios without any prior information (Mask Kneser-Ney vs. Whole Kneser-Ney).

password guessing scenarios. We employ each approach to pre-generate 10^6 passwords for each template. For existing sequence models (i.e., Markov [43], OMEN [20], and Backoff [43] and PCFG [67]), we first enable them to pre-generate 10^9 passwords in descending order of probability. Then we filter out passwords that conform templates PTs from these guesses, and match V^{PT} to obtain the guessing success rates. We call this approach the default mode of a sequence model (i.e., “-d” in Table VII). While these models are not inherently designed for mask guessing, they can be better applied to this attack scenario with proper adjustments. We call this approach the mask mode of a sequence model (i.e., “-m” in Table VII).

Algorithm 1 illustrates how to adapt the generation mechanism of sequence models from traditional whole-password guessing (default mode) to template-based guessing (mask mode). The core idea is to perform depth-first search (DFS) under template and probability constraints. For example, given a password template $123*5*$, each trained sequence model M (including our PassSeq and Kneser-Ney) provides a character distribution for each wildcard position $*$. The constrained DFS procedure for generating guesses that naturally match a given template (e.g., $123*5*$) generally consists of four steps. First, we set a minimum probability threshold $\mathcal{T}$ and initialize DFS from left to right. Second, if the character at the current position is known (e.g., 5 in $123*5*$), we directly append it to the current prefix $s(=1234)$. Third, if the position corresponds to a wildcard $*$, we enumerate candidate characters c whose conditional probability satisfies $\Pr_M(s||c) > \mathcal{T}$. Fourth, we recursively repeat the above process until the end of the password template is reached.

Results. Table VII reports the comparison results across four template classes. We observe that, when adapted to the mask guessing scenario, traditional statistical models (e.g., Markov [43] and Backoff [43]) achieve substantial improve-

Algorithm 1: PW generation for mask guessing.

Input: Template PT, Model, Threshold $\mathcal{T}$, CharSet.

Output: Generate passwords pw that match the given template.

```

1 DFS(pw, p, i) : /* p is the probability of password prefix; i
   is the i-th position of pw; i < len(pw); i initialize to 0 */
2 if i = (PT.length - 1) then
3   return pw
4 if PT[i] = * then
5   for c in CharSet do
6     pw += c
7     if Model.calculate(pw) < T then
8       continue
9   else
10    p = Model.calculate(pw)
11    DFS(pw, p, i + 1) /* Search from the next position. */
12 else
13   pw += PT[i] /* Add the unmasked character directly. */
14   if Model.calculate(pw) < T then
15     return
16   else
17     p = Model.calculate(pw)
18     DFS(pw, p, i + 1)

```

ments in guessing capability (see the difference between the “-d” and “-m” columns). Surprisingly, Backoff-m outperforms state-of-the-art neural models that specifically tailored for mask guessing (including CPG [52], PassBERT [71], and RankGuess-MASK [72]), across all template categories. This demonstrates that employing statistical smoothing to estimate character-level probabilities is a promising and effective strategy for tackling the mask password guessing problem.

In particular, our Kneser-Ney model (KN-m) achieves the highest guessing success rates among all baseline models (excluding PassSeq) across the four template classes. It is especially effective in the *rare* category, where data sparsity

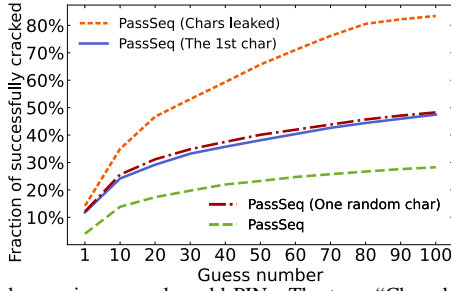


Fig. 5. Mask guessing on real-world PINs. The term “Chars leaked” means that the $\mathcal{A}_{\text{stat}}$ attacker knows the digit composition of the target victim’s PIN but does not know the exact positions of the digits. PassSeq is trained on 50% Amitay-4-digit dataset, and is tested on the rest 50%.

at higher-order n -gram levels poses significant modeling challenges and smoothing becomes crucial. While slightly behind KN-m and Backoff-m on the first three higher-frequency template categories, our PassSeq demonstrates its greatest strength in the most challenging *super-rare* class. Specifically, PassSeq achieves a success rate of 0.4879, outperforming the strongest neural baseline RankGuess-MASK [72] (0.3147) by a relative margin of 55.1%, and our enhanced Backoff-m model (0.4094) by 19.2%. These results highlight PassSeq’s robustness in addressing extremely sparse password templates, where most prior models fail to generalize effectively.

To qualitatively compare bidirectional encoding (PassSeq) with n -gram smoothing (Kneser-Ney), we examine the templates from Table VII for which PassSeq achieves fewer average guesses than Kneser-Ney. Results show a general trend: PassSeq performs better than Kneser-Ney when masked positions are irregularly distributed, widely separated, or disrupt the continuity of local character sequences (e.g., `*am**bon*`; see more examples in Appendix Table XX). Such templates are generally longer and contain multiple fragmented masked blocks, which introduce long-range dependencies that n -gram smoothing cannot effectively capture. PassSeq, benefiting from bidirectional encoding, aggregates contextual information across distant unmasked segments and remains effective even when local substrings are heavily interrupted.

Further, we have set up some general scenarios to show the impact of exposing/masking *different numbers of characters* on password security. Results (see Table X in Appendix B) show that disclosing just one *additional* character can *double* the attacker’s guessing success rate (e.g., a 121.9% increase within 10^5 guesses on the COMB dataset using PassSeq).

Since PassSeq, Kneser-Ney, and our improved Backoff baseline (Backoff-m) consistently rank among the top three in guessing success rates and outperform state-of-the-art mask-guessing models such as RankGuess-MASK [72], PassBERT [71], and CPG [52], we mainly employ these three models in the following experiments to *quantitatively evaluate the threat posed by different classes of mask-guessing attackers*.

C. Distributional mask guessing scenarios

Considering that some side-channel attacks (e.g., thermal residue [4]) may not reveal any particular password character with 100% certainty, we set up two mask attack scenarios in

which the attacker acquires statistical or compositional side-channel information about the target password.

Thermal attackers. To model probabilistic character leakage in mask guessing attacks, we construct per-position character disclosure probabilities inspired by empirical observations from prior thermal-residue attack studies. More specifically, we adopt a six-position recognizability profile, $[0.17, 0.20, 0.25, 0.33, 0.50, 1.00]$, inspired by the monotonically increasing position-wise decoding probabilities illustrated in Fig. 11 and Table I of ThermoSecure [6], which reflect progressively stronger confidence in character ordering as the decoding proceeds. We use these values as heuristic position-wise disclosure priors for character leakage. We then linearly interpolate these values to any password length L to obtain a position-dependent probability sequence q_i . Each q_i is clipped to $[0, 1]$ and perturbed with Gaussian noise $\mathcal{N}(0, 0.05^2)$ to account for real-world variability (e.g., surface material). Experimental results show that the guessing success rates of disclosing characters (with thermal-residue-based leakage probabilities) are 2.47-33.67 times higher than those of not disclosing any password character (i.e., the Whole Kneser-Ney lines in Fig. 4). Fig. 14 (in Appendix) additionally considers other leakage probabilities that follow the increasing-probability trend observed in thermal-residue attacks [4], [6], and the conclusions still hold.

Smudge attackers. In practice, smudges in keypads (see Fig. 1) can reveal the target victim’s character composition information (without order). We evaluate the impact of this threat on real-world PINs using the Amitay-4-digit dataset, which consists of 204,432 human-chosen 4-digit PINs collected by Daniel Amitay with the iOS application “Big Brother Camera Security” in 2011. Specifically, we assume that $\mathcal{A}_{\text{stat}}$ knows the digits composing the target 4-digit PIN through keystroke traces but does not know their exact positions (which is also aligned with the experimental setup in [14]). To model this, we modify the PassSeq model to limit the search space to the known digits of the target user when generating guesses. We then calculate the guessing success rate in this scenario and compare it to the scenario where the attacker has no prior knowledge. Fig. 5 shows that the success rate of a *single guess* by $\mathcal{A}_{\text{stat}}$ is 255% higher than when the attacker has no prior knowledge. Within 10 guesses, the success rate of the attacker $\mathcal{A}_{\text{stat}}$ increases by 152%, reaching 35%.

D. Contextual prior mask guessing scenarios

In the above, we have only considered the threat posed by mask attackers who possess partial information about the target victim’s password, without leveraging the target victim’s PII or previously leaked passwords. Yet, in reality, a large fraction of users (e.g., 37%-51% in [64]) build passwords with their own PII, and tend to reuse their existing passwords (e.g., 20%-59% of users [17], [64] directly reuse or simply modify their existing passwords). Moreover, the large-scale data breaches (e.g., 2.9 billion individuals’ personal information has been leaked on the dark web [2]) provide attackers with ample material for training models. Thus, it is necessary and practical

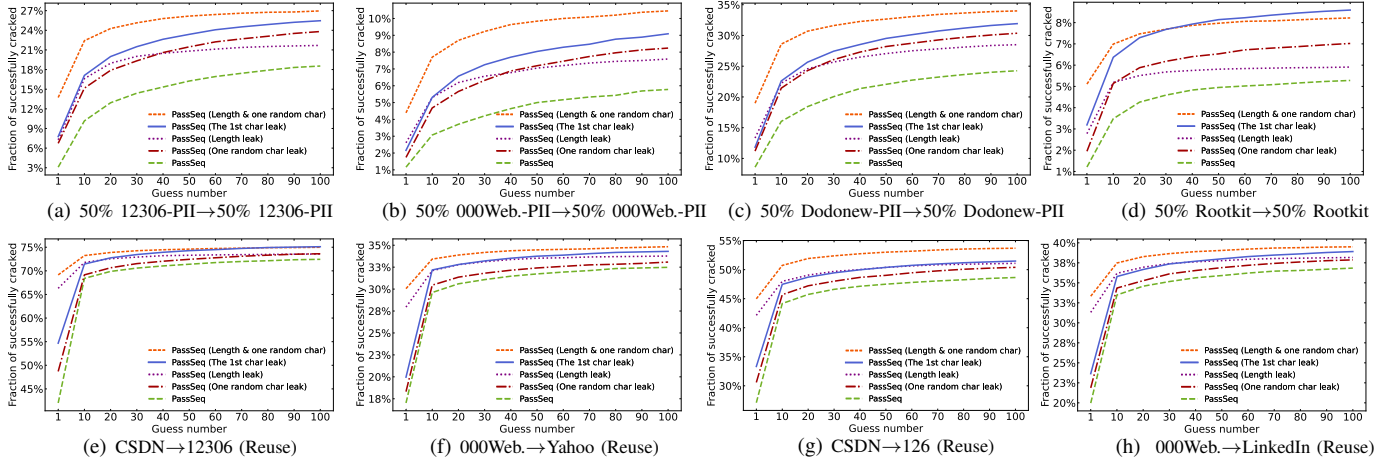


Fig. 6. Mask guessing scenarios with PII/existing leaked password (i.e., the $\mathcal{A}_{\text{context}}$ attacker). The training and testing experimental setup is recommended by [64]. On both Chinese and English datasets, disclosing password length and one password character significantly improves the attacker’s guessing success rate. In Figs. 6(a)-6(d), with a single guess, the leak of password length and one character *doubles* the attacker’s success rates (see Table XI for the confidence statistics of the success-rate improvement), increasing by at least $165\% \pm 24\%$ (95% CI: 124%-223%). See Table XIV for specific success-rate figures.

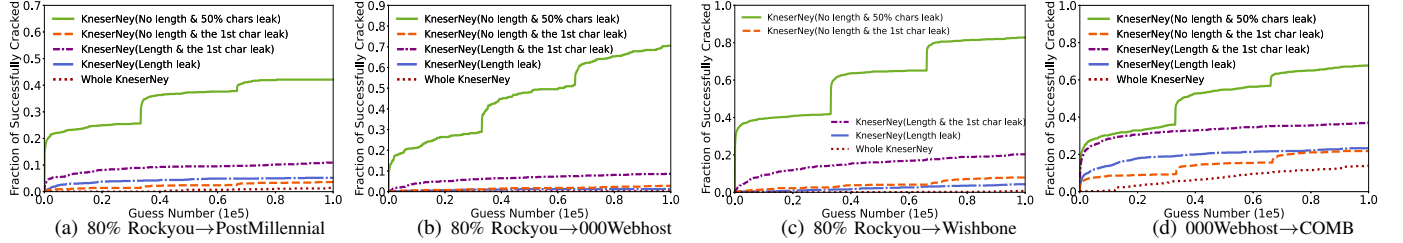


Fig. 7. Length-uncertain mask guessing with/without one-character leakage (the $\mathcal{A}_{\text{length}}$ and $\mathcal{A}_{\text{stat}}$ attackers). We consider five attacker configurations with different prior knowledge: (1) no prior information (whole-password guessing), (2) unknown length with probabilistic character leakage ($p=0.5$), (3) unknown length with one leaked character, (4) known length without character leakage, and (5) known length with one leaked character. These scenarios collectively represent the upper (100% leakage) and lower bounds (0% leakage) of leaking length and/or one character. Our results show that at 10^5 guesses, leaking the first character (with unknown password length) at least doubles the attacker’s success rate, achieving improvements from 1.57-14.50 times.

to evaluate the threat posed by a more powerful mask attacker knowing the target victim’s PII or existing passwords. Tables V and VI summarize the basic information of our datasets.

PII-based attackers. For the PII-based mask attack scenario, we train our PassSeq on datasets containing PII (see Table V). More specifically, the training input for our PassSeq is the character sequence of the personal information sequence (e.g., name|username|email|birthday), and the output is the corresponding password sequence. When generating guesses, we input the target victim’s PII sequence, using beam search to generate the corresponding guesses. Figs. 6(a) and 6(b) show that by merely knowing the victim’s password length, the guessing success rates of the $\mathcal{A}_{\text{context}}$ attacker can increase by 63.48%-72.64% within 10 guesses, and this value *doubles* (+102.17%-129.32%) when one additional password character is leaked. He et al. [26] demonstrate that inconsistent PII-masking policies across apps can be stitched together to reconstruct users’ PII, which can then be used to abuse password-recovery to reset/recover passwords. Our focus is complementary but different: rather than using PII itself as a first or secondary form of authentication and password recovery, we quantify the incremental risk when PII-based targeted guessing is further combined with partial leaks such as password length and/or characters.

Reuse-based attackers. For the reuse-based mask attack scenario, we train the PassSeq model by taking a user’s existing password as input and their new password as output, following a training paradigm similar to Pass2Pass [49] and PointerGuess [69]. We model a realistic attacker $\mathcal{A}_{\text{context}}$ who knows both the existing password and the target password’s length and/or one character. Here, our PassSeq serves as a representative model to evaluate the *additional* risks introduced by mask guessing in password reuse settings. Importantly, our goal is not to compare PassSeq with state-of-the-art reuse-based models (indeed, it achieves guessing success rates comparable to PointerGuess [69]). To validate our findings, we also apply the same evaluation to PointerGuess (see Figs. 13(a)-13(d) in Appendix), and observe consistent trends. Results (see Figs. 6(e)-6(h)) show that with only one guess, the guessing success rates of $\mathcal{A}_{\text{context}}$ increase by 62.30%-73.21% when both the password length and one character are disclosed, with the single-guess success rate reaching up to 69% (see Fig. 6(e)). All this reveals that mask guessing with users’ PII or existing passwords is a damaging threat that needs more attention.

E. Uncertain-length mask guessing scenarios

To quantify the impact of *password length* information on mask guessing scenarios, we assume that the attacker only

TABLE VIII
PERFORMANCE COMPARISON ACROSS DIFFERENT MODELS[†].

Model	Training time	Model size	Password generation speed	
			Trawling (pw/s)	Mask (pw/s)
Kneser-Ney	00:05:45	3.0 GB	284,000	2,042,222
PassSeq	84:38:58	23.0 MB	610	151
RankGuess-MASK [72]	16:08:11	2.52 MB	/	30
PassBERT [71]	22:00:08	35.9 MB	/	320
CPG [52]	26:54:04	6.1 MB	/	677
Backoff [‡]	00:01:55	270.0 MB	297,600	2,195,121

[†] CPU: Xeon(R) Gold 6226R 2.9GHz; GPU: RTX 3090 (80% Rockyou).

[‡] We optimize the original Backoff [43] implementation to achieve theoretical optimal speed, resulting in several times speedup (see Appendix G for details).

knows partial character information of the target password but without knowing the exact password length. In this case, an intuitive strategy is to prioritize the selection of the top lengths counted in the training set as the priority choices for constructing the templates. More specifically, we choose the top-3 lengths from the training set to set three different lengths for each test template. Here, the top-3 lengths we select should satisfy the condition of being greater than or equal to the longest position of the leaked character in the target password that we know. For example, if we know that the character at the eighth position of the target password is *d*, then the selected top lengths should be ≥ 8 .

Consider a toy example where the target password is *p@ssword*. A password template with masks is generated by masking each position with a 50% probability, and without loss of generality, assume the template **@s*w**d* is generated. Then, from the leaked dataset (i.e., the training set), $\mathcal{A}_{\text{length}}$ (who knows characters at positions 2, 4, 5 and 8, but *not* the total length) selects the top-3 common lengths (≥ 8 ; e.g., 8, 9, and 10) and constructs three templates (i.e., **@s*w**d*, **@s*w***d** and **@s*w****d***). $\mathcal{A}_{\text{length}}$ then generates an equal number of guesses for each template (1/3 of the total) to execute the attack. Similar to the PII/reuse-based guessing scenarios, we additionally consider three combined leakage cases in trawling scenarios: (i) length-only leakage, (ii) single-character leakage without password length, and (iii) simultaneous leakage of both length and one character. As shown in Figs. 7(a)-7(d), leaking the password length and the first character can substantially increase the attacker’s guessing success rate. Notably, even when the password length is unknown, leaking only the first character is sufficient to at least *double* the attacker’s success rate, with improvements of 1.57-14.50 times within 10^5 guesses.

F. Performance evaluation

Table VIII shows that adapting existing statistical models such as Backoff [43] and Kneser-Ney (KN) to mask guessing scenarios significantly improves their generation speeds. In contrast, our PassSeq has slower generation speed in mask guessing scenarios (151 pw/s vs. 610 pw/s), primarily because mask generation relies on sequential token-by-token decoding with limited parallelism, compared to the more batch-friendly trawling mode. However, the slower generation speed of PassSeq is acceptable in practice. Mask guessing is typically employed in targeted online attacks, where the performance bottleneck lies in network latency and server-side throttling

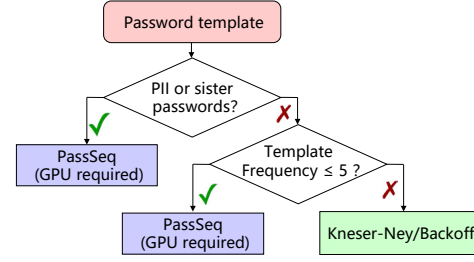


Fig. 8. Decision chart for mask password model selection. It is recommended to use our PassSeq model when the target victim’s PII or sister passwords are available. For password templates with a frequency ≤ 5 , PassSeq is preferred; otherwise, Kneser-Ney or Backoff is recommended.

mechanisms [21], rather than in local computational cost. As such, PassSeq remains practical for real-world deployment.

It is worth noting that, besides achieving the highest overall guessing success rates in typical mask scenarios (see Table VII), our KN offers fast training and generation speeds without requiring GPU support (see Table VIII), making it suitable for *on-site* training without the need to save/maintain model files. This property is quite desirable. Meanwhile, our Seq2Seq-based PassSeq framework supports a wider range of mask guessing scenarios (including PII- and reuse-based attacks) with strong generalization ability. Importantly, it performs best under *super-rare* template types (see Table VII), where statistical models generally struggle due to data sparsity. **Recommended usage guidelines.** Table VIII shows that PassSeq outperforms KN and Backoff on templates with frequencies in the range of [1, 5], whereas its success rate is lower than KN when the template frequency increases to [10, 15]. To further examine this trend, we evaluate PassSeq, KN, and Backoff on templates with frequencies between [5, 10], where their guessing success rates are 93.93%, 98.02%, and 92.86%, respectively. Considering all these results alongside Fig. 6, we recommend using PassSeq for templates with frequencies ≤ 5 and for targeted scenarios involving PII or previously leaked passwords, while KN/Backoff are better suited for higher-frequency templates or deployments that prioritize fast, GPU-free training and inference (see Fig. 8).

Overhead of MLE. Note that neural password models (e.g., our PassSeq) require a significant amount of time to generate a large number of guesses (e.g., $>10^7$) that *naturally conform to the given password templates*. For example, it takes PassSeq 18.5 hours to explicitly generate 10^7 guesses conforming a *single* given template with four or five masks. Fortunately, our MLE can reduce the time cost of guess number estimation (for any password under the given template) to minutes. Given a template containing three masks (e.g., **w*rt*12*), there are a total of 94^3 unique passwords under this template. We employ MLE to estimate the guess number required to crack a target password (e.g., *qwErTy12*) under this template. The MLE parameters are set as follows: the number of samples $N=10^5$, and the number of repetitions $t=4$. The average time consumed by MLE to estimate the guess number for any password under this given template (i.e., **w*rt*12*) is 107s (using PassSeq). Notably, the time consumption is mainly in the sampling and probability calculation process for a given template. Once this

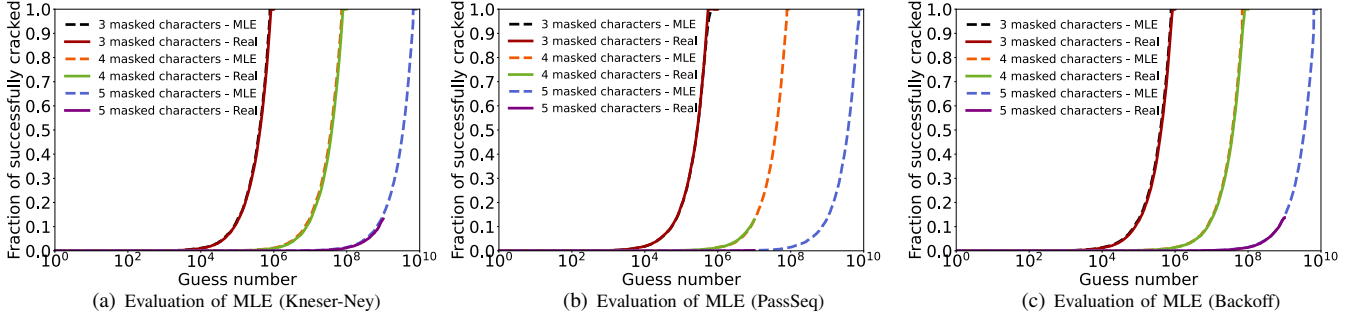


Fig. 9. Comparison between the actual password generation results (e.g., 3 masked characters-Real) and the MLE estimates (e.g., 3 masked characters-MLE) for Kneser-Ney, PassSeq, and Backoff [43] under three template types with different numbers of masks ($n=3, 4, 5$). Results show that the MLE estimates closely match the empirical generation outcomes, with the maximum deviation in guessing success rates $<1\%$.

process is completed, the time for estimating the guess number of any 94^4 passwords under this password template is $<0.1s$.

Evaluation of MLE. To empirically evaluate the effectiveness of our proposed MLE estimator, we categorize password templates by the number of masks they contain, namely three, four, and five masks, which correspond to increasingly large guessing spaces (e.g., a four-mask template corresponds to a space of size 94^4). These three categories represent the most frequently occurring template types in typical mask guessing scenarios and are therefore selected as representative cases. For each category (a total of 30 templates are tested), we conduct random sampling with sample sizes of 10^4 , 10^5 , and 10^6 , and repeat the sampling process four times, using 10,000 randomly selected passwords as the test set under the given template. In these settings, the MLE results exhibit stable convergence. We further applied PassSeq, Kneser-Ney, and Backoff [43] to perform both actual password generation and MLE-based estimation across the three template types. Considering that Kneser-Ney and Backoff generate guesses more efficiently than PassSeq, their actual generation sizes are larger than PassSeq (10^9 vs. 10^7). Fig. 9 shows that the empirical guessing success rates closely match the MLE estimates, with a maximum deviation of less than 1%, demonstrating the estimator’s accuracy and robustness across different models and templates.

G. A further validation of PassSeq

To evaluate the practical effectiveness of PassSeq under realistic side-channel attack scenarios, we integrate it into keystroke acoustic inference attacks (see Fig. 10). Following prior work [23], [73], [74], we collect keystroke audio from 11 keyboard units across commonly used brands and models at realistic eavesdropping distances (17-30 cm), involving 16 diverse participants. Among them, nine distinct keyboards are used for same-device/same-user evaluations, each providing 100 keystroke audio samples recorded by individual participants entering 6-digit PINs sampled from the top-100 most frequent PINs in the Taobao dataset. To improve robustness and capture subtle acoustic variations, we additionally record 800 samples (top-800 PINs) on a 2020 MacBook Air entered by six different participants. Moreover, two identical keyboards of the same brand and model are typed by two users to enable realistic cross-user evaluations. Each same-device dataset is



Fig. 10. A further validation: mask guessing with keystroke acoustic information leaked from different types keyboards. (The experimental setup we used for keystroke collection is the same as [23], [73], [74]).

split into training and test sets with an 80:20 ratio, and each is used to train an independent acoustic inference model.

For acoustic inference, we start from the open-source CNN-based implementation of [23] and further enhance it following [3] by integrating CNNs, random forest classifiers, and MFCC-based feature extraction into a unified ensemble framework. This improves the single-key inference accuracy from about 40% to 94% under same-device in-domain splits. This ensemble model is adopted as our final acoustic backbone.

Besides same-device evaluations, we further conduct cross-device experiments to examine generalization. Specifically, we train and test on different keyboards from the same brand (and vice versa). We investigate two strategies for integrating PassSeq: (1) *Template-driven integration* (oracle distribution), where keystroke acoustic posteriors are used to construct candidate password templates (e.g., $1*2**3$) that guide PassSeq’s generation; and (2) *Probability-fusion integration* (no-oracle distribution), where posterior probabilities from the acoustic and password models are directly combined using normalized weights. Based on empirical tuning, we assign higher weights (0.8-0.9) to the acoustic model under same-device settings, while under cross-device settings, a lower acoustic weight (0.1-0.2) achieves better results. Moreover, we have explored a complementary approach in cross-device settings, where acoustic posteriors are used to reorder PassSeq’s beam-search outputs. Its effectiveness is comparable to that of character-level probability fusion, with neither approach consistently outperforming the other.

TABLE IX
GUESSING SUCCESS RATES OF THE ACOUSTIC, PASSSEQ, AND INTEGRATED MODELS ACROSS DIFFERENT KEYBOARDS[†].

Keyboard type / Cross-device Setting	Acoustic model [3], [23]			PassSeq model			Acoustic+PassSeq model v1 [‡]			Acoustic+PassSeq model v2 [‡]		
	Top-1	Top-10	Top-100	Top-1	Top-10	Top-100	Top-1	Top-10	Top-100	Top-1	Top-10	Top-100
HP EliteBook	30.00%	60.00%	80.00%	0.00%	0.00%	0.00%	45.00%	75.00%	85.00%	15.00%	55.00%	80.00%
Logitech K220	35.00%	50.00%	55.00%	0.00%	5.00%	20.00%	85.00%	90.00%	90.00%	35.00%	55.00%	85.00%
Apple Keyboard	40.00%	75.00%	95.00%	0.00%	0.00%	0.00%	40.00%	85.00%	90.00%	40.00%	75.00%	85.00%
Varmilo	75.00%	90.00%	95.00%	0.00%	0.00%	0.00%	55.00%	95.00%	95.00%	75.00%	95.00%	95.00%
Mech. KB	25.00%	30.00%	70.00%	0.00%	0.00%	0.00%	25.00%	55.00%	80.00%	30.00%	55.00%	85.00%
Dell KB216T	15.00%	25.00%	25.00%	10.00%	15.00%	25.00%	55.00%	65.00%	70.00%	15.00%	40.00%	60.00%
Dell KB212-B	15.00%	20.00%	25.00%	0.00%	10.00%	20.00%	25.00%	35.00%	45.00%	10.00%	30.00%	65.00%
Lenovo	15.00%	15.00%	60.00%	0.00%	0.00%	20.00%	30.00%	40.00%	45.00%	30.00%	55.00%	75.00%
MacBook Air 25	15.00%	25.00%	40.00%	0.00%	5.00%	25.00%	40.00%	55.00%	55.00%	20.00%	50.00%	65.00%
MacBook Air 20	18.75%	39.38%	53.12%	0.62%	2.50%	6.88%	28.75%	47.50%	60.00%	18.12%	28.75%	36.88%
Dell KB216T → Dell KB212-B	0.00%	0.00%	0.00%	1.00%	3.00%	24.00%	0.00%	0.00%	0.00%	1.00%	4.00%	24.00%
Dell KB212-B → Dell KB216T	0.00%	0.00%	0.00%	2.00%	9.00%	21.00%	0.00%	0.00%	0.00%	1.00%	2.00%	28.00%
MacBook Air 25 → MacBook Air 20	0.62%	0.62%	0.62%	0.25%	0.88%	3.14%	0.62%	0.62%	0.62%	0.25%	1.13%	3.39%
MacBook Air 20 → MacBook Air 25	0.00%	0.00%	0.00%	1.00%	4.00%	18.00%	0.00%	0.00%	0.00%	1.00%	4.04%	21.21%
Dell KB216T _A → Dell KB216T _B	7.00%	10.00%	21.00%	0.00%	3.00%	17.00%	14.00%	21.00%	25.00%	7.00%	15.00%	17.00%
Dell KB216T _B → Dell KB216T _A	7.00%	8.00%	15.00%	1.00%	3.00%	14.00%	9.00%	14.00%	17.00%	4.00%	6.00%	14.00%

[†] All percentages represent Top- k success rates under each model. Each within-device test sets $N = 20$ (except MacBook Air 20 with $N = 160$). Cells with dark gray and light gray backgrounds denote the highest and second-highest values (per keyboard and Top- k) across the four models, respectively.

[‡] v1 means using the acoustic model’s predicted character probabilities to construct candidate password templates (e.g., $1*2**3$), corresponding to the oracle distribution. v2 directly adds the acoustic character probabilities to those produced by PassSeq, forming the no-oracle distribution.

The overall template construction procedure is driven by per-keystroke acoustic posteriors. For each segmented keystroke, we extract three uncertainty-related statistics from the posterior distribution: (1) the top-1 probability, (2) the top- k candidates ($k=3$), and (3) the entropy to quantify prediction uncertainty. Based on these statistics, each position is either fixed to its most likely digit or masked, according to a set of entropy- and confidence-thresholding rules. Specifically, positions with high confidence (low entropy or high top-1 probability) are fixed, while uncertain positions are replaced by a mask token. To further enrich template diversity, we additionally generate single-position-masked variants from the top acoustic sequence prediction, as well as global Top- N masking templates that retain only the most reliable positions. All generated templates are then evaluated by a posterior-guided scoring function considering per-position confidence, consistency with the base acoustic prediction, and the proportion of fixed characters. The top-ranked templates (e.g., top-5 per acoustic candidate) are finally selected to guide the subsequent PassSeq-based password generation.

For PassSeq (implemented using a lightweight Seq2Seq architecture with two unidirectional LSTM layers), we use 6-digit numeric PINs from the Dodonew dataset as the training set. Guided by the constructed templates, the trained PassSeq generates ranked password guesses under both the template-driven and probability-fusion integration strategies.

Overall, integrating the acoustic model with the password model significantly improves guessing success rates across most keyboards (see Table IX). Specifically, the oracle distribution integration (v1) achieves substantial gains: within 10 guesses, it generally outperforms the acoustic model by 5.6%-166.7%, and the PassSeq model by 3-18 times. The no-oracle distribution integration (v2) also outperforms the acoustic-only model in 13/16 cases, though to a lesser extent. Notably, in cross-device scenarios, the acoustic-only and oracle models perform poorly due to strong inter-device variability. While the

v2 integration still provides measurable improvements, most of its gains stem from the password model itself. In contrast, in cross-user scenarios (two users typing on two identical keyboards, i.e., Dell KB216T_A and Dell KB216T_B), the oracle integration significantly outperforms the acoustic baseline by 12.5%-110% within 10 guesses.

V. PRACTICAL DEFENSES AGAINST MASK GUESSING

Employ fully static masking and minimize UI feedback granularity. At CCS’24, Hu et al. [33] found that dynamic masking, widely adopted on mobile browsers, provides little usability benefit compared with fully static masking. Together with our findings that partial character exposure (particularly the first character; see Fig. 6) significantly boosts attackers’ guessing success rates, we recommend that mobile platforms also adopt fully static masking (already widely deployed on desktop systems [33]) and avoid any deterministic per-keystroke reveals or position-dependent hints. Such uniform masking effectively mitigates visual side-channel leakage without significantly sacrificing usability.

Blind password length across UI, transport, and application layers. Prior work by Harsha et al. [24] demonstrates that >84% of Alexa Top-100 websites are vulnerable to leak exact password length through request sizes even under TLS. Consistent with their findings, our evaluation indicates that leaking only the password length can increase an attacker’s guessing success rate by 1.57 times (see Fig. 7). We therefore recommend end-to-end length blinding: using fixed-width input fields and placeholder elements on the client, disabling length-threshold hints, and padding both requests and error responses so that their observable sizes remain uniform.

Apply multi-dimensional adaptive rate limiting. Current throttling implementations remain inconsistent, effectively allowing thousands of online guesses per month (e.g., Amazon, which enforces the strictest lockout policy among the Alexa Top-10 websites, allows 3,600 monthly guessing attempts; see

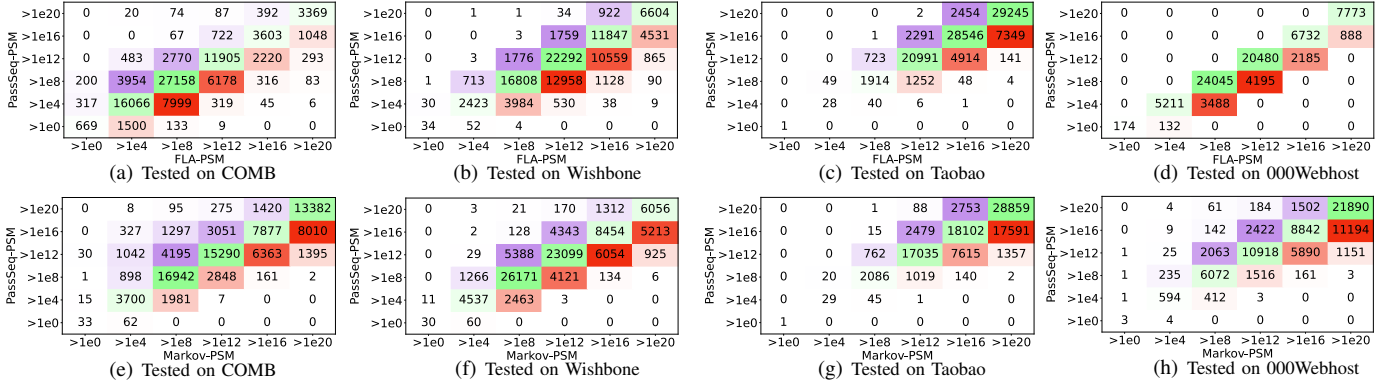


Fig. 11. We compare PassSeq-PSM with FLA-PSM [46]/Markov-PSM [15] using the unsafe error metric across four different password datasets. PassSeq-PSM significantly outperforms FLA-PSM [46]. For example, in Figure 11(a), our PassSeq-PSM evaluates 7,999 passwords as being guessable between 10^4 and 10^8 guesses, while the FLA-PSM [46] rates them as requiring more than 10^8 guesses. Overestimates of strength are indicated in shades of red, underestimates in purple, and accurate estimates in green. The intensity of the color increases with the number of passwords in each cell.

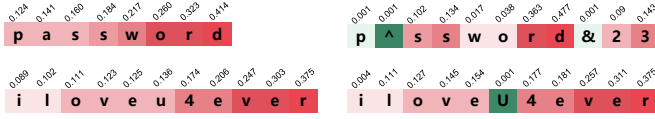


Fig. 12. Character-level feedback examples for two password pairs. Different colors indicate the estimated character security levels: Red (insecure)→Green (secure). PassSeq automatically identifies weak characters by assigning high values to them. We keep four significant digits, where the same score corresponds to lighter colors, indicating higher security.

Table 7 in [62]). This tolerance suggests that our proposed password models remain effective across both online and offline scenarios, since our experimental configurations span guess budgets from 10^0 - 10^6 , fully covering the scale of typical online limits. To counter such structured, mask-aware attacks, we recommend behavior-driven adaptive throttling: dynamically adjusting thresholds based on account history, device fingerprint, and the presence of mask-like guessing attempts (e.g., fixed-length sequences, position-constrained character classes, or leakage-conditioned templates).

Mitigate acoustic leakage and password-level inference.

Our results show that the integration of per-keystroke acoustic leakage and password models amplifies the attacker’s capability in two ways: keystroke sounds expose character-level cues, while password models capture structural and semantic characteristics. We thus provide two complementary mitigation directions. On the signal side, device- or application-level masking, such as mixing background sounds (which can be generated on the input device and does not impede users from entering passwords or performing tasks effectively [41]), has been shown to substantially reduce key-specific recognition cues. On the user side, adjusting typing behavior before entering sensitive information, and creating long, non-semantic passwords limits the gain attackers obtain from both acoustic features and password model-driven inference.

Practical applications for protection. By employing the Monte Carlo method [18], PassSeq can serve as a password strength meter (PSM), providing a security evaluation of passwords from an attacker’s perspective. We evaluate PassSeq-PSM using the unsafe error metric (i.e., the rate at which a

PSM incorrectly labels weak or easily guessable passwords as strong, a widely adopted measure in prior work [46], [70]). We compare PassSeq-PSM with FLA-PSM [46] and Markov-PSM [15] on 100,000 randomly sampled passwords from the 000Webhost, Wishbone, Taobao, and COMB datasets. Fig. 11 shows that our PassSeq-PSM has significantly fewer unsafe errors than both baselines, demonstrating improved reliability in password strength evaluation.

Besides overall strength evaluation, PassSeq-PSM enables fine-grained analysis by introducing a Markov Random Field (MRF) to model the security impact between characters, thereby quantifying each character’s contribution to the total password strength. Specifically, each character in the password represents a state in MRF, and each state is connected by two undirected edges with weights. We denote a given password as $pw=c_1c_2\dots c_i\dots c_j\dots c_l$, the security impact of the character c_i on c_j as SI_i^j , and the uncertainty of c_i (i.e., the negative impact of c_i on password security) as SI_i . Then, we can calculate the edge weight $w(i, j)$ from c_i to c_j in MRF as

$$w(i, j) = -\text{Softmax}(SI_i^j) \cdot \log_2[\text{Softmax}(SI_i^j)]. \quad (11)$$

If c_i and c_j are not connected, $w(i, j) = 0$. We denote the sum of weights of edges starting from c_i as $w(i)$, then

$$w(i) = \sum_{1 \leq t \leq l, t \neq i} w(i, t) = SI_i. \quad (12)$$

Based on MRF, the strength of pw evaluated by our PassSeq-PSM is calculated as

$$\text{Strength}(pw) = \frac{\sum_{1 \leq i \leq l} w(i)}{Z}, \quad (13)$$

where Z is used to normalize the probability, and it is set according to the Hammersley-Clifford theorem.

In this way, PassSeq-PSM can calculate and feedback the security impact/contribution of a character on both the whole password (i.e., $w(i)$) and the other characters (i.e., $w(i, j)$). A toy example is shown in Figure 12.

VI. CONCLUSION

We have presented the first systematic study of how attackers can substantially amplify password guessing effectiveness by exploiting *partial password information* (e.g., some characters and/or password length) combined with side-channel priors, PII, and sister passwords. Extensive experiments with 15 real-world datasets demonstrate the effectiveness and wide applicability of our proposed PassSeq and Kneser-Ney. Our results highlight the damaging threat of mask guessing, especially when combined with keystroke inference attacks or PIIs. We believe that the new algorithms and insights into mask guessing threats can help re-evaluate existing password practices and inform future password security research.

ACKNOWLEDGEMENT

The authors are deeply grateful to the shepherd and anonymous reviewers for their invaluable comments. Ding Wang is the corresponding author. This research was in part supported by the National Natural Science Foundation of China under Grants Nos. 62502238 and 62572259, and by the China Postdoctoral Science Foundation under Grant No. 2025M771596, and by the Tianjin Natural Science Foundation Project under Grant No. 25JCJQC00230.

ETHICAL CONSIDERATIONS

Our work is inherently dual-use. While it advances the understanding of password security under partial-information leakage (e.g., side-channel priors) and helps defenders quantify real-world risk, the same techniques could be misused to facilitate more efficient password guessing if deployed irresponsibly. We explicitly acknowledge this risk and adopt concrete mitigation measures to ensure that the defensive and societal benefits outweigh the potential for misuse.

Specifically, our study uses publicly available password datasets that were leaked through past breaches and have been widely used in prior password research (e.g., [43], [46], [52], [64], [71], [72]). However, besides passwords, our study incorporates associated PIIs (e.g., usernames or emails) in the context of targeted guessing scenarios. We recognize that affected individuals did not provide informed consent for research use of their data and that incorporating PII raises greater ethical concerns than password-only studies. While these datasets are already publicly accessible, their continued use still warrants careful ethical consideration.

To minimize potential harm to data subjects, we adopt strict data-handling safeguards. We do not redistribute any raw breached datasets, PII, or unprocessed side-channel traces, and our study does not introduce any additional exposure beyond what already exists in the public domain. All datasets are stored and processed on offline machines not connected to the Internet. Individual accounts are treated as strictly confidential, and all reported results are aggregated and anonymized, ensuring that no personally identifiable information or individual account details can be inferred. Sensitive attributes, including PII, are deleted immediately after the completion of analysis. All acoustic data are collected with explicit participant consent under controlled conditions.

REFERENCES

- [1] *Global Visual Hacking Experimental Study: Analysis*, Aug. 2016, <https://multimedia.3m.com/mws/media/1254232O/global-visual-hacking-experiment-study-summary.pdf>.
- [2] *Your Social Security Number Has Probably Been Stolen – Here’s What To Do*, Aug 2024, <https://tech.co/news/social-security-number-stolen-a-dvice>.
- [3] *Keystroke Acoustic Recognition System*, April 2025, https://github.com/s2cr3t/Keystroke_Acoustic_Recognition_System.
- [4] Y. Abdelrahman, M. Khamis, S. Schneegass, and F. Alt, “Stay cool! understanding thermal attacks on mobile-based user authentication,” in *Proc. ACM CHI 2017*, pp. 3751–3763.
- [5] P. Alexander, *John the Ripper community build (1.9.0-bleeding jumbo)*, 2025, <http://www.openwall.com/john/>.
- [6] N. Alotaibi, J. Williamson, and M. Khamis, “Thermosecure: Investigating the effectiveness of AI-driven thermal attacks on commonly used computer keyboards,” *ACM Trans. on Priv. and Secur.*, vol. 26, no. 12, pp. 1–24, 2023.
- [7] A. J. Aviv, J. T. Davin, F. Wolf, and R. Kuber, “Towards baselines for shoulder surfing on mobile authentication,” in *Proc. ACSAC 2017*, pp. 486–498.
- [8] A. J. Aviv, K. L. Gibson, E. Mossop, M. Blaze, and J. M. Smith, “Smudge attacks on smartphone touch screens,” in *Proc. USENIX WOOT 2010*, pp. 1–7.
- [9] J. Bai, B. Liu, and L. Song, “I know your keyboard input: A robust keystroke eavesdropper based-on acoustic signals,” in *Proc. ACM MM 2021*, pp. 1239–1247.
- [10] K. S. Balagani, M. Cardaioli, S. Ceconello, M. Conti, and G. Tsudik, “We can hear your PIN drop: An acoustic side-channel attack on ATM PIN pads,” in *Proc. ESORICS 2022*, pp. 633–652.
- [11] F. Binbeshr, M. L. M. Kiah, L. Y. Por, and A. A. Zaidan, “A systematic review of pin-entry methods resistant to shoulder-surfing attacks,” *Comput. Secur.*, vol. 101, no. 102116, 2021.
- [12] J. Bonneau, C. Herley, P. van Oorschot, and F. Stajano, “Passwords and the evolution of imperfect authentication,” *Commun. ACM*, vol. 58, no. 7, pp. 78–87, 2015.
- [13] J. Bonneau, C. Herley, P. C. Van Oorschot, and F. Stajano, “The quest to replace passwords: A framework for comparative evaluation of web authentication schemes,” in *Proc. IEEE S&P 2012*, pp. 553–567.
- [14] M. Cardaioli, M. Conti, K. S. Balagani, and P. Gasti, “Your PIN sounds good! augmentation of PIN guessing strategies via audio leakage,” in *Proc. ESORICS 2020*, pp. 720–735.
- [15] C. Castelluccia, M. Dürmuth, and D. Perito, “Adaptive password-strength meters from markov models,” in *Proc. NDSS 2012*, pp. 1–14.
- [16] S. Cha, S. Kwag, H. Kim, and J. H. Huh, “Boosting the guessing attack performance on android lock patterns with smudge attacks,” in *AsiaCCS. ACM*, 2017, pp. 313–326.
- [17] A. Das, J. Bonneau, M. Caesar, N. Borisov, and X. Wang, “The tangled web of password reuse,” in *Proc. NDSS 2014*, pp. 1–15.
- [18] M. Dell’Amico and M. Filippone, “Monte carlo strength evaluation: Fast and reliable password checking,” in *Proc. ACM CCS 2015*, pp. 158–169.
- [19] Q. Dong, C. Jia, F. Duan, and D. Wang, “RLS-PSM: A robust and accurate password strength meter based on reuse, leet and separation,” *IEEE Trans. Inf. Forensics Secur.*, vol. 16, pp. 4988–5002, 2021.
- [20] M. Dürmuth, F. Angelstorf, C. Castelluccia, D. Perito, and C. Abdelber, “OMEN: faster password guessing using an ordered markov enumerator,” in *Proc. ESSoS 2015*, pp. 119–132.
- [21] M. Dürmuth, D. Freeman, S. Jain, B. Biggio, and G. Giacinto, “Who are you? A statistical approach to measuring user authenticity,” in *Proc. NDSS 2016*, pp. 1–15.
- [22] M. Eiband, M. Khamis, E. von Zezschwitz, H. Hussmann, and F. Alt, “Understanding shoulder surfing in the wild: Stories from users and observers,” in *Proc. CHI 2017*, pp. 4254–4265.
- [23] J. Harrison, E. Toreini, and M. Mehrzad, “A practical deep learning-based acoustic side channel attack on keyboards,” in *Proc. EuroS&P Workshops*, 2023, pp. 270–280.
- [24] B. Harsha, R. Morton, J. Blocki, J. A. Springer, and M. Dark, “Bicycle attacks considered harmful: Quantifying the damage of widespread password length leakage,” *Comput. Secur.*, vol. 100, p. 102068, 2021.
- [25] Hashcat Project, “Hashcat,” <https://hashcat.net/hashcat/>.
- [26] F. He, Y. Jia, J. Zhao, Y. Fang, J. Wang, M. Feng, P. Liu, and Y. Zhang, “Maginot line: Assessing a new cross-app threat to pii-as-factor authentication in chinese mobile apps,” in *Proc. NDSS 2024*.

- [27] K. Heafield, I. Pouzyrevsky, J. H. Clark, and P. Koehn, “Scalable modified kneser-ney language model estimation,” in *Proc. ACL 2013*, 2013, pp. 690–696.
- [28] C. Herley and P. Van Oorschot, “A research agenda acknowledging the persistence of passwords,” *IEEE Secur. Priv.*, vol. 10, pp. 28–36, 2012.
- [29] B. Hitaj, P. Gasti, G. Ateniese, and F. Pérez-Cruz, “PassGAN: A deep learning approach for password guessing,” in *Proc. ACNS 2019*.
- [30] S. Hochreiter and J. Schmidhuber, “Long short-term memory,” *Neural Comput.*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [31] B. Honan, *Visual Data Security White Paper*, July 2012, <https://multimedia.3m.com/mws/media/950026O/secure-white-paper.pdf>.
- [32] S. Houshmand, S. Aggarwal, and R. Flood, “Next gen PCFG password cracking,” *IEEE Trans. Inf. Forensics Secur.*, vol. 10, no. 8, pp. 1776–1791, 2015.
- [33] Y. Hu, S. Alroomi, S. Sahin, and F. Li, “Unmasking the security and usability of password masking,” in *Proc. CCS 2024*, pp. 4241–4255.
- [34] Z. Huang, D. Wang, and Y. Zou, “Prob-hashcat: Accelerating probabilistic password guessing with hashcat by hundreds of times,” in *Proc. RAID 2024*, pp. 674–692.
- [35] A. Juels and R. L. Rivest, “Honeywords: Making password-cracking detectable,” in *Proc. ACM CCS 2013*, pp. 145–160.
- [36] S. Katz, “Estimation of probabilities from sparse data for the language model component of a speech recognizer,” *IEEE Trans. Acoust. Speech Signal Process.*, vol. 35, no. 3, pp. 400–401, 1987.
- [37] P. G. Kelley, S. Komanduri, M. L. Mazurek, R. Shay, T. Vidas, L. Bauer, N. Christin, L. F. Cranor, and J. Lopez, “Guess again (and again and again): Measuring password strength by simulating password-cracking algorithms,” in *Proc. IEEE S&P 2012*, pp. 523–537.
- [38] R. Kneser and H. Ney, “Improved backing-off for M-gram language modeling,” in *Proc. IEEE ICASSP 1995*, pp. 181–184.
- [39] T. Kwon, S. Shin, and S. Na, “Covert attentional shoulder surfing: Human adversaries are more powerful than expected,” *IEEE Trans. Syst. Man Cybern. Syst.*, vol. 44, no. 6, pp. 716–727, 2014.
- [40] E. Liu, A. Nakanishi, M. Golla, D. Cash, and B. Ur, “Reasoning analytically about password-cracking software,” in *Proc. IEEE S&P 2019*, pp. 380–397.
- [41] A. Lobanova, J. He, N. Zhu, A. Nazir, and X. Ma, “Machine-learning-based adaptive keyboard sound masking against acoustic channel attacks,” in *Proc. CCNS 2024*, pp. 271–277.
- [42] T. Luong, H. Pham, and C. D. Manning, “Effective approaches to attention-based neural machine translation,” in *Proc. ACL EMNLP 2015*, pp. 1412–1421.
- [43] J. Ma, W. Yang, M. Luo, and N. Li, “A study of probabilistic password models,” in *Proc. IEEE S&P 2014*, pp. 689–704.
- [44] R. Manning, *2025 Global State of Authentication survey: A world of difference in cybersecurity habits*, Sept. 2025, <https://www.yubico.com/blog/2025-global-state-of-authentication-survey-a-world-of-difference-in-cybersecurity-habits/>.
- [45] D. Marques, I. Muslukhov, T. J. Guerreiro, L. Carriço, and K. Beznosov, “Snooping on mobile phones: Prevalence and trends,” in *Proc. SOUPS 2016*, pp. 159–174.
- [46] W. Melicher, B. Ur, S. Komanduri, L. Bauer, N. Christin, and L. F. Cranor, “Fast, lean and accurate: Modeling password guessability using neural networks,” in *Proc. USENIX SEC 2017*, pp. 175–191.
- [47] R. Morris and K. Thompson, “Password security: A case history,” *Commun. ACM*, vol. 22, no. 11, pp. 594–597, 1979.
- [48] A. Narayanan and V. Shmatikov, “Fast dictionary attacks on passwords using time-space tradeoff,” in *Proc. ACM CCS 2005*, pp. 364–372.
- [49] B. Pal, T. Daniel, R. Chatterjee, and T. Ristenpart, “Beyond credential stuffing: Password similarity models using neural networks,” in *Proc. IEEE S&P 2019*, pp. 417–434.
- [50] D. Pasquini, G. Ateniese, and M. Bernaschi, “Interpretable probabilistic password strength meters via deep learning,” in *Proc. ESORICS 2020*.
- [51] D. Pasquini, G. Ateniese, and C. Troncoso, “Universal neural-cracking-machines: Self-configurable password models from auxiliary data,” in *Proc. IEEE S&P 2024*, pp. 36–54.
- [52] D. Pasquini, A. Gangwal, G. Ateniese, M. Bernaschi, and M. Conti, “Improving password guessing via representation learning,” in *Proc. IEEE S&P 2021*, pp. 1382–1399.
- [53] J. Rando, F. Pérez-Cruz, and B. Hitaj, “Passgpt: Password modeling and (guided) generation with large language models,” in *Proc. ESORICS 2023*, pp. 164–183.
- [54] E. Shareghi, M. Petri, G. Haffari, and T. Cohn, “Fast, small and exact: Infinite-order language modelling with compressed suffix trees,” *Trans. Assoc. Comput. Linguistics*, vol. 4, pp. 477–490, 2016.
- [55] X. Su, X. Zhu, Y. Li, Y. Li, C. Chen, and P. J. E. Verissimo, “Pagpassgpt: Pattern guided password guessing via generative pretrained transformer,” in *Proc. DSN 2024*, pp. 429–442.
- [56] I. Sutskever, O. Vinyals, and Q. V. Le, “Sequence to sequence learning with neural networks,” in *Proc. NIPS 2014*, pp. 3104–3112.
- [57] C. Team, *New Study Underscores Slow Adoption of Multifactor Authentication By Global SMBs*, Nov. 2024, <https://cyberreadinessinstitute.org/news-and-events/new-study-underscores-slow-adoption-of-multifactor-authentication/>.
- [58] L. Van der Maaten and G. Hinton, “Visualizing data using t-SNE,” *J. Mach. Learn. Res.*, vol. 9, no. 11, 2008.
- [59] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, and I. Polosukhin, “Attention is all you need,” in *Proc. NIPS 2017*, pp. 5998–6008.
- [60] R. Veras, C. Collins, and J. Thorpe, “On semantic patterns of passwords and their security impact,” in *Proc. NDSS 2014*.
- [61] C. Wang, J. Zhang, M. Xu, H. Zhang, and W. Han, “#segments: A dominant factor of password security to resist against data-driven guessing,” *Comput. Secur.*, vol. 121, no. 102848, 2022.
- [62] D. Wang, X. Shan, Q. Dong, Y. Shen, and C. Jia, “No single silver bullet: Measuring the accuracy of password strength meters,” in *Proc. USENIX SEC 2023*, pp. 947–964.
- [63] D. Wang, P. Wang, D. He, and Y. Tian, “Birthday, name and bifacial-security: Understanding passwords of Chinese web users,” in *Proc. USENIX SEC 2019*, pp. 1537–1555.
- [64] D. Wang, Z. Zhang, P. Wang, J. Yan, and X. Huang, “Targeted online password guessing: An underestimated threat,” in *Proc. ACM CCS 2016*, pp. 1242–1254.
- [65] D. Wang, Y. Zou, Q. Dong, Y. Song, and X. Huang, “How to attack and generate honeywords,” in *Proc. IEEE S&P 2022*, pp. 489–506.
- [66] D. Wang, Y. Zou, Z. Zhang, and K. Xiu, “Password guessing using random forest,” in *Proc. USENIX SEC 2023*, pp. 965–982.
- [67] M. Weir, S. Aggarwal, B. De Medeiros, and B. Glodek, “Password cracking using probabilistic context-free grammars,” in *Proc. IEEE S&P 2009*, pp. 391–405.
- [68] D. Wheeler, “zxcvbn: Low-budget password strength estimation,” in *Proc. USENIX SEC 2016*, pp. 157–173.
- [69] K. Xiu and D. Wang, “Pointerguess: Targeted password guessing model using pointer mechanism,” in *Proc. USENIX SEC 2024*, pp. 5555–5572.
- [70] M. Xu, C. Wang, J. Yu, J. Zhang, K. Zhang, and W. Han, “Chunk-level password guessing: Towards modeling refined password composition representations,” in *Proc. ACM CCS 2021*, pp. 5–20.
- [71] M. Xu, J. Yu, X. Zhang, C. Wang, S. Zhang, H. Wu, and W. Han, “Improving real-world password guessing attacks via bi-directional transformers,” in *Proc. USENIX SEC 2023*, pp. 1001–1018.
- [72] T. Yang and D. Wang, “Rankguess: Password guessing using adversarial ranking,” in *Proc. S&P 2025*, pp. 682–700.
- [73] J. Yu, L. Lu, Y. Chen, Y. Zhu, and L. Kong, “An indirect eavesdropping attack of keystrokes on touch screen through acoustic sensing,” *IEEE Trans. Mob. Comput.*, vol. 20, no. 2, pp. 337–351, 2021.
- [74] T. Zhu, Q. Ma, S. Zhang, and Y. Liu, “Context-free attacks using keyboard acoustic emanations,” in *Proc. ACM CCS 2014*.
- [75] V. Zimmermann, “From the quest to replace passwords towards supporting secure and usable password creation,” Ph.D. dissertation, Technical University of Darmstadt, Germany, 2021.
- [76] Y. Zou, M. An, and D. Wang, “Password guessing using large language models,” in *Proc. USENIX SEC 2025*, pp. 7799–7818.

APPENDIX

A. A toy example of the Kneser-Ney smoothing

Consider a character-level 4-gram model trained on the following password corpus: {Passwd, Pastel, Pastime, Misdeal, Session1–5, Designs, Signs, Messages, Stressful, Tossed}.

Assume we aim to predict the next character given the prefix Pas. The substring Pasd does not appear in the training set, i.e., $C(\text{Pasd}) = 0$, but related substrings like Pass and Past

TABLE X
COMPARISON OF GUESSING SUCCESS RATES OF DIFFERENT APPROACHES (RANDOMLY MASK 3, 4, 5 CHARACTERS IN A PASSWORD).[†]

Experimental setup		Masked character #=3			Masked character #=4				Masked character #=5				
Test set	Guesser	10 ³	10 ⁴	10 ⁵	10 ⁴	10 ⁵	10 ⁶	10 ⁷	10 ⁵	10 ⁶	10 ⁷	10 ⁸	10 ⁹
000Webhost	Markov [43]	56.87%	73.66%	83.96%	38.71%	58.79%	75.68%	89.74%	26.59%	43.38%	58.63%	77.07%	93.29%
	Backoff [43]	57.23%	73.68%	84.16%	42.43%	59.12%	75.73%	88.18%	27.16%	46.68%	61.09%	77.42%	94.94%
	FLA [46]	57.12%	73.23%	83.67%	43.16%	59.05%	77.69%	91.01%	27.20%	45.56%	59.11%	76.13%	93.34%
	Hashcat [25]	43.12%	66.92%	78.39%	30.27%	43.86%	56.92%	70.97%	15.01%	31.39%	43.40%	53.83%	67.49%
	Kneser-Ney	59.32%	78.94%	87.19%	46.38%	60.76%	78.31%	93.94%	30.81%	47.45%	59.84%	78.01%	96.25%
	PassSeq	57.17%	76.22%	89.04%	45.63%	61.68%	79.51%	92.59%	29.24%	48.52%	61.15%	79.54%	98.33%
COMB	Markov [43]	59.50%	84.93%	95.80%	35.14%	56.11%	73.72%	89.47%	24.19%	47.78%	65.90%	82.92%	94.64%
	Backoff [43]	60.48%	85.65%	95.91%	43.24%	58.23%	76.64%	90.11%	27.15%	48.13%	67.85%	84.76%	96.91%
	FLA [46]	59.22%	84.89%	95.76%	42.17%	59.31%	76.21%	90.33%	27.58%	51.84%	67.02%	86.31%	96.12%
	Hashcat [25]	52.22%	69.21%	80.11%	27.56%	53.82%	70.17%	80.18%	19.07%	37.88%	47.29%	56.62%	63.59%
	Kneser-Ney	64.32%	85.32%	95.96%	45.89%	60.73%	77.43%	93.21%	29.24%	51.46%	69.99%	87.73%	98.05%
	PassSeq	62.19%	85.14%	95.82%	47.93%	62.64%	77.01%	95.97%	28.20%	53.87%	67.04%	88.19%	97.63%
Taobao	Markov [43]	69.94%	93.11%	99.01%	40.28%	60.22%	73.72%	85.72%	30.92%	52.54%	74.17%	89.81%	96.53%
	Backoff [43]	70.14%	94.97%	99.02%	46.32%	60.22%	76.64%	86.24%	32.51%	53.03%	74.19%	89.46%	97.96%
	FLA [46]	65.34%	89.61%	94.44%	48.11%	62.32%	76.21%	85.33%	33.20%	51.33%	74.01%	91.36%	97.78%
	Hashcat [25]	45.32%	65.32%	82.56%	30.62%	45.27%	70.17%	75.27%	21.16%	34.37%	47.11%	59.58%	64.10%
	Kneser-Ney	70.32%	95.20%	99.32%	47.31%	65.81%	76.72%	88.24%	35.21%	56.93%	78.76%	93.45%	98.24%
	PassSeq	70.19%	94.97%	99.15%	47.17%	66.07%	78.02%	89.15%	37.24%	55.06%	77.12%	93.10%	98.42%
Wishbone	Markov [43]	65.31%	89.08%	97.56%	39.33%	62.72%	83.23%	95.28%	31.86%	47.66%	64.83%	83.32%	96.78%
	Backoff [43]	66.11%	89.32%	98.45%	41.36%	64.28%	85.56%	95.15%	33.28%	48.65%	67.96%	85.29%	97.09%
	FLA [46]	65.31%	86.12%	96.32%	42.18%	64.38%	84.02%	94.96%	35.34%	49.70%	65.12%	85.24%	96.43%
	Hashcat [25]	44.15%	65.32%	75.18%	31.62%	54.34%	74.07%	86.67%	24.97%	32.21%	45.08%	56.67%	63.59%
	Kneser-Ney	67.78%	88.51%	98.94%	45.25%	65.66%	86.65%	97.64%	39.18%	51.35%	67.82%	89.58%	98.26%
	PassSeq	65.98%	90.87%	98.75%	47.68%	66.68%	87.94%	97.14%	38.10%	50.37%	68.38%	89.02%	98.51%

[†]A **bold** value (success rate) means that it is the *highest* one in each test dataset. Masking different character # corresponds to different upper limit of guesses.

appear in total three times. Under a Markov password model [43] using MLE, the conditional probability is:

$$\Pr(d | Pas) = \frac{C(Pasd)}{C(Pas)} = \frac{0}{3} = 0. \quad (14)$$

This zero probability excludes the possibility of generating $Pasd$, even though d may be a reasonable continuation.

Backoff smoothing [43] alleviates this issue by backing off to lower-order statistics. For example, the bigram sd appears once (in *Misdeal*), while s appears 17 times overall. Thus, $\Pr(d|s) = \frac{1}{17} \approx 0.0588$. Since Pas appears three times and is followed by two distinct characters (i.e., s and t), the backoff weight is:

$$\alpha(Pas) = \frac{0.75 \cdot 2}{3} = 0.5. \quad (15)$$

Here, 0.75 is a commonly used discount constant in Katz backoff smoothing. Hence, the smoothed probability becomes:

$$\Pr_{\text{Backoff}}(d | Pas) = 0.5 \cdot 0.0588 \approx 0.0294. \quad (16)$$

While backoff avoids zero probabilities, it does not account for the diversity of contexts in which d appears.

In contrast, Kneser-Ney smoothing [38] estimates the probability of a character not just based on frequency but also on the number of distinct contexts in which it occurs. In our training set, d appears after three different characters: w in *Passwd*, s in *Misdeal*, s in *Misdeal*, and e in *Tossed*. The continuation count for d is thus 3. Given that there are 44 distinct bigrams in total, the continuation probability is:

$$\Pr_{\text{cont}}(d) = \frac{3}{44} \approx 0.0682. \quad (17)$$

Assuming a similar backoff weight $\gamma(Pas) = 0.5$ (for comparison with backoff), KN smoothing gives:

$$\Pr_{\text{KN}}(d | Pas) \approx \gamma(Pas) \cdot \Pr_{\text{cont}}(d) = 0.5 \cdot 0.0682 \approx 0.0341. \quad (18)$$

Compared to $\Pr_{\text{Backoff}}(d|Pas)=0.029$, Kneser-Ney assigns a higher probability to d due to its higher contextual diversity. **Summary.** Compared to the traditional Markov and Backoff password models [43], Kneser-Ney offers two main advantages in password modeling: (1) it assigns non-zero, context-aware probabilities to previously unseen character sequences, and (2) it distributes probabilities more effectively under sparse data by incorporating the diversity of preceding contexts through continuation counts. These strengths make it especially suitable for mask password guessing, where character combinations are generally sparse and irregular.

B. Impact of leaked/masked character numbers

In practice, the position and number of leaked password characters obtained by the attacker vary according to the actual situation. Without loss of generality, we randomly *mask* different numbers of characters in a password, and Table X summarizes the results of masking 3, 4, and 5 characters. Here we give the concrete values at $[10^n, 10^{n+1}, \dots, 10^{2n-1}]$ guesses, where n stands for the number of masked characters. Results show that reducing the number of masked characters (i.e., increasing the number of leaked characters) significantly boosts the guessing success rates across all datasets and models. For example, the guessing success rates of our PassSeq increase by 75.0%-121.9% when reducing the number of masked characters from 5 to 4 within 10^5 guesses.

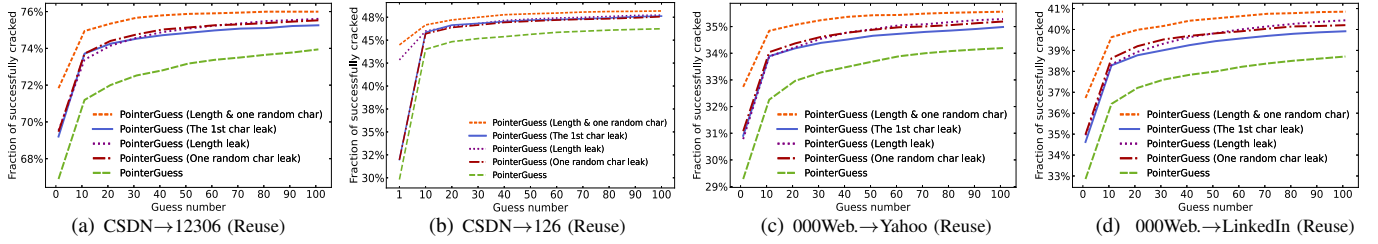


Fig. 13. We employ the state-of-the-art password reuse model PointerGuess [69] to evaluate mask password guessing given access to users’ existing (sister) passwords (i.e., $\mathcal{A}_{\text{context}}$). The results consistently show that disclosing partial password information, such as the password length and/or a single character, substantially boosts the attacker’s guessing success rates across all evaluated datasets. In particular, the improvement is most pronounced at the very early stage of the attack (e.g., the first guess), underscoring the risk posed by mask password guessing in reuse-based attack scenarios.

TABLE XI

CONFIDENCE STATISTICS OF SUCCESS-RATE IMPROVEMENTS (ONE CHARACTER+LENGTH LEAKAGE VS. BASELINE) UNDER DIFFERENT SCENARIOS AND GUESS BUDGETS[†].

Scenarios	Top-1 (Mean±Std, 95% CI)	Top-10 (Mean±Std, 95% CI)	Top-100 (Mean±Std, 95% CI)
(a) 50% 12306 → 50% 12306	3.4150 ± 0.6312 (2.4405, 4.9938)	1.2331 ± 0.2207 (0.9289, 1.7605)	0.4893 ± 0.0665 (0.3752, 0.6356)
(b) 50% 000Webhost → 50% 000Webhost	2.7961 ± 1.1789 (1.3225, 5.9781)	1.5555 ± 0.3672 (0.9554, 2.3875)	0.8666 ± 0.1425 (0.5727, 1.1932)
(c) 50% Dodonew → 50% Dodonew	1.6532 ± 0.2439 (1.2416, 2.2259)	0.8900 ± 0.1036 (0.7060, 1.0787)	0.4029 ± 0.0500 (0.2974, 0.4872)
(d) 50% Rootkit → 50% Rootkit	3.4139 ± 1.0793 (1.3747, 5.9562)	0.9676 ± 0.2131 (0.6201, 1.3814)	0.5276 ± 0.1115 (0.2428, 0.7561)

[†]All values in the table are relative improvement ratios. Each result is obtained by performing 30 rounds of sampling with replacement over the entire test set, with 1,000 target users per round. Cells highlighted in gray correspond to scenarios where the success rates doubled compared to the baseline.

TABLE XII

CONFIDENCE STATISTICS OF SUCCESS-RATE IMPROVEMENTS (ONE CHARACTER LEAKAGE VS. BASELINE) UNDER DIFFERENT SCENARIOS AND GUESS BUDGETS[†].

Scenarios	Top-1 (Mean±Std, 95% CI)	Top-10 (Mean±Std, 95% CI)	Top-100 (Mean±Std, 95% CI)
(a) 50% 12306 → 50% 12306	1.5686 ± 0.3802 (1.0935, 2.3839)	0.6397 ± 0.0970 (0.4797, 0.9004)	0.4116 ± 0.0588 (0.3055, 0.5555)
(b) 50% 000Webhost → 50% 000Webhost	0.8257 ± 0.3770 (0.2493, 2.0628)	0.7437 ± 0.1993 (0.3841, 1.1790)	0.6866 ± 0.1242 (0.4342, 0.9149)
(c) 50% Dodonew → 50% Dodonew	0.5948 ± 0.0895 (0.4249, 0.7677)	0.4692 ± 0.0576 (0.3505, 0.5927)	0.3295 ± 0.0439 (0.2200, 0.3842)
(d) 50% Rootkit → 50% Rootkit	1.7479 ± 0.5563 (0.6727, 2.9475)	0.8241 ± 0.1805 (0.5724, 1.2675)	0.6344 ± 0.1163 (0.4491, 0.8819)

[†]All values in the table are relative improvement ratios. Each result is obtained by performing 30 rounds of sampling with replacement over the entire test set, with 1,000 target users per round. Cells highlighted in gray correspond to scenarios where the success rates doubled compared to the baseline.

TABLE XIII

CONFIDENCE STATISTICS OF SUCCESS-RATE IMPROVEMENTS (LENGTH LEAKAGE VS. BASELINE) UNDER DIFFERENT SCENARIOS AND GUESS BUDGETS[†].

Scenarios	Top-1 (Mean±Std, 95% CI)	Top-10 (Mean±Std, 95% CI)	Top-100 (Mean±Std, 95% CI)
(a) 50% 12306 → 50% 12306	1.4074 ± 0.3149 (0.9449, 1.9943)	0.6426 ± 0.1227 (0.4540, 0.9132)	0.1846 ± 0.0366 (0.1159, 0.2728)
(b) 50% 000Webhost → 50% 000Webhost	1.1918 ± 0.4791 (0.4247, 2.3224)	0.7173 ± 0.1939 (0.3496, 1.1143)	0.3514 ± 0.0931 (0.2174, 0.5476)
(c) 50% Dodonew → 50% Dodonew	0.7015 ± 0.1049 (0.5001, 0.9501)	0.4564 ± 0.0559 (0.3601, 0.5591)	0.1823 ± 0.0411 (0.1074, 0.2638)
(d) 50% Rootkit → 50% Rootkit	1.3689 ± 0.5381 (0.5865, 2.6935)	0.4497 ± 0.1111 (0.2425, 0.6734)	0.1115 ± 0.0518 (0.0192, 0.2466)

[†]All values in the table are relative improvement ratios. Each result is obtained by performing 30 rounds of sampling with replacement over the entire test set, with 1,000 target users per round. Cells highlighted in gray correspond to scenarios where the success rates doubled compared to the baseline.

C. Template masking methods and their empirical evaluation

In this section, we present three masking methods for constructing password templates and evaluate their impacts on PassSeq’s guessing success rates. The three methods (i.e., random masking, hotword masking, and character-to-character impact masking) represent progressively refined strategies that incorporate increasingly fine-grained security considerations.

Random masking method. In the random masking method, each character in a password is independently masked with probability p . This produces templates with different levels of structural exposure and allows PassSeq to observe a diverse range of template patterns during training. To determine an effective value for p , we evaluate $p \in \{0.2, 0.5, 0.8\}$ on the Rockyou and Sohu datasets using an 80%/20% train/test split. As shown in Table XV, within

5×10^8 guesses, the guessing success rate follows the order random-0.5 > random-0.8 > random-0.2. When p is too large, too few characters remain exposed and the templates provide insufficient structural guidance. When p is too small, the templates become overly rigid and limit the diversity of generated guesses. A moderate masking level, namely $p=0.5$, offers a balance between structural constraints and generative flexibility, which leads to a slightly higher success rate.

Hotword masking method. Intuitively, PassSeq readily captures high-frequency substrings (i.e., hotwords) that appear in password datasets. Consider $\text{pw}=\text{iloveyou}$: a template such as $\text{PT}_{\text{pw}}^{\text{ideal}} = *\text{love}***$ provides strong structural cues, enabling PassSeq to reconstruct the original password more reliably than with a template like $*1**\text{eyo}*$. Motivated by this observation, we distinguish between hotword and non-

TABLE XIV
MASK GUESSING SUCCESS RATES (%) AT 1, 10, AND 100 GUESSES UNDER EIGHT PII/PASSWORD-REUSE SCENARIOS (FIG. 6).

Scenario	Leakage information	Top-1	Top-10	Top-100
(a) 50% 12306→12306	No leakage	3.09%	10.16%	18.54%
	One random character	6.73%	15.16%	23.81%
	First character	7.76%	17.11%	25.46%
	Password length	7.27%	16.61%	21.69%
	Length + one character	13.73%	22.41%	26.96%
(b) 50% 000Web→000Web	No leakage	1.16%	3.07%	5.78%
	One random character	1.74%	4.68%	8.24%
	First character	2.16%	5.31%	9.09%
	Password length	2.61%	5.30%	7.59%
	Length + one character	4.39%	7.70%	10.45%
(c) 50% Dodonew→Dodonew	No leakage	8.58%	16.02%	24.26%
	One random character	11.19%	21.43%	30.36%
	First-character	11.82%	22.60%	31.92%
	Password length	13.32%	22.26%	28.50%
	Length + one character	18.98%	28.59%	33.99%
(d) 50% Rootkit→Rootkit	No leakage	1.21%	3.48%	5.28%
	One random character	1.95%	5.14%	7.02%
	First character	3.19%	6.37%	8.59%
	Password length	2.78%	5.19%	5.91%
	Length + one character	5.11%	6.99%	8.22%
(e) CSDN→12306 (Reuse)	No leakage	42.02%	68.41%	72.44%
	One random character	48.70%	69.15%	73.62%
	First character	54.68%	71.44%	75.11%
	Password length	66.30%	71.78%	73.56%
	Length + one character	69.14%	73.23%	74.99%
(f) 000Web→Yahoo (Reuse)	No leakage	17.02%	29.61%	32.48%
	One random character	18.31%	30.45%	33.08%
	First character	19.96%	32.18%	34.30%
	Password length	27.95%	32.10%	33.77%
	Length + one character	30.05%	33.42%	34.83%
(g) CSDN→126 (Reuse)	No leakage	27.09%	44.17%	48.66%
	One random character	30.58%	45.69%	50.40%
	First character	33.33%	47.47%	51.48%
	Password length	42.14%	47.99%	51.09%
	Length + one character	44.98%	50.73%	53.69%
(h) 000Web→LinkedIn (Reuse)	No leakage	19.97%	33.49%	36.84%
	One random character	21.87%	34.35%	37.87%
	First character	23.65%	35.79%	38.91%
	Password length	31.31%	36.19%	38.17%
	Length + one character	33.33%	37.48%	39.50%

TABLE XV
EVALUATION OF THE RANDOM MASKING METHOD.

Random masking	Guessing success rate (%) of PassSeq			
	Rockyou (training size:test size=8:2)	Sohu (training size:test size=8:2)		
probability	1×10^8 guesses	5×10^8 guesses	1×10^8 guesses	5×10^8 guesses
0.2	23.22	49.01	23.38	45.65
0.5	25.31	52.24	23.22	48.32
0.8	23.36	50.11	23.21	46.54

hotword characters and treat each hotword as an indivisible unit during masking, meaning that all characters within a hotword are considered to contribute equally to its security impact. To extract hotwords, we count every substring p_s of length at least three, remove those with frequency below $\mathcal{T}_1$, and adjust their counts using

$$C(p_s)^{\text{new}} = C(p_s)^{\text{old}} - \sum_{c \in \Sigma} [C(c + p_s)^{\text{old}} + C(p_s + c)^{\text{old}}]. \quad (19)$$

Substrings with $C(p_s)^{\text{new}} > \mathcal{T}_2$ are identified as hotwords. In

TABLE XVI
EVALUATION OF THE HOTWORD MASKING METHOD.

Hotword masking probability	Guessing success rate (%) of PassSeq			
	Rockyou (training size:test size=8:2)	Sohu (training size:test size=8:2)		
	1×10^8 guesses	5×10^8 guesses	1×10^8 guesses	5×10^8 guesses
0.0	21.37	50.01	21.21	45.33
0.2	23.68	53.88	23.34	49.15
0.8	23.57	51.35	23.01	46.42
1.0	22.19	49.85	22.98	45.01

TABLE XVII
COMPARISON OF THE THREE MASKING METHODS.

Optimal masking methods	Guessing success rate (%) of PassSeq			
	Rockyou (training size:test size=8:2)	Sohu (training size:test size=8:2)		
	1×10^8 guesses	5×10^8 guesses	1×10^8 guesses	5×10^8 guesses
Random-0.5	25.31	52.24	23.22	48.32
Hotword-0.2	23.68	53.88	23.34	49.15
C2C-impact [†]	24.70	54.35	23.38	50.35

[†]C2C-impact is short for Character-to-character impact.

our experiments, we empirically set $(\mathcal{T}_1, \mathcal{T}_2) = (100, 5)$ for hotword identification. Each hotword is masked with probability p , while non-hotword characters are masked with $1-p$.

We evaluate $p \in \{0.0, 0.2, 0.8, 1.0\}$ and obtain the guessing success-rate ranking hotword-0.2>hotword-0.8>hotword-0.0>hotword-1.0, as shown in Table XVI. This behavior can be explained as follows. When too many hotwords are masked, the model loses common password-building components that are crucial for realistic generation. When no hotwords are masked, the resulting templates become overly constrained. A small masking rate, specifically $p=0.2$, preserves essential structural information while still allowing sufficient variation for PassSeq to generalize.

Character-to-character impact masking method. Beyond frequency-based structure, the presence of one character may influence the likelihood of other characters in the same password. To capture this dependency, we define the security impact of c_i on c_j in a password $\text{pw} = c_1 \dots c_i \dots c_j \dots c_l$. After masking positions i and j , we compute

$$SI_i^j = P(A_j | A_i) = \frac{P(A_i \cap A_j)}{P(A_i)}, \quad (20)$$

where the probabilities are obtained from PassSeq:

$$P(A_i) = \sum_{\tilde{c}_j \in \Sigma} P(c_1 \dots c_i \dots \tilde{c}_j \dots c_l | \text{PT}_{\text{pw}}), \quad (21)$$

$$P(A_i \cap A_j) = P(c_1 \dots c_i \dots c_j \dots c_l | \text{PT}_{\text{pw}}). \quad (22)$$

We normalize $\{SI_i^j\}_{j \neq i}$ using Softmax and compute the entropy

$$SI_i = - \sum_{j \neq i} \text{Softmax}(SI_i^j) \log_2[\text{Softmax}(SI_i^j)]. \quad (23)$$

Characters are ranked in ascending order of SI_i , and the least informative $\lceil l/2 \rceil$ characters are masked.

Comparison and overall insights. Table XVIII presents representative decoding outcomes obtained using our three masking methods on several example passwords from the Rockyou dataset, including the highly popular password (rank 4)

TABLE XVIII
PASSWORD TEMPLATES AND THE TOP-20 DECODING RESULTS OBTAINED BY APPLYING OUR THREE MASKING METHODS.[†]

Example password	password			superman40			shoes25		
Masking method	Random	Hotword	C2C-impact	Random	Hotword	C2C-impact	Random	Hotword	C2C-impact
Password template	pa*s***d	p*s*s***d	p***w*rd	****rman*0	su***m*n40	s**er***n*0	s*oe*2*	sh***s*5	s*o***5
Rank	Top-20 decoding results								
1	password	password	password	superman10	superman40	superman10	stoe123	shans15	stone25
2	passw0rd	passw0rd	passw0rd	loverman10	sunnymon40	superman20	snoe123	shans05	smokey5
3	passwind	passwind	passward	laverman10	sunnyman40	superman00	shoe123	shoes25	stone15
4	passwood	passwood	passwird	superman20	sunshman40	supermin10	sloe123	shans95	stone05
5	passwird	passwird	passwurd	hotarman10	succemon40	superman90	skoe123	shans25	sport15
6	password	password	passwlrđ	angerman10	sundamon40	superman80	scoet23	shars15	shoes25
7	passwold	passwold	pattword	interman10	supermin40	supermine0	scoet22	shans75	spooky5
8	passwond	passwold	papsword	superman00	summeman40	superman30	syoe123	shans45	stong15
9	passw0od	passw0od	pa55word	daverman10	summymon40	superson10	scoet20	shels95	stone85
10	passwwod	passwwod	partword	madarman10	succeman40	supermon10	scoet21	shels15	scooty5
11	passwurd	passwurd	pastword	liverman10	sundaman40	saverman10	scoe123	shans85	shone15
12	passiond	passiond	pardword	superman40	sunnamon40	superman60	stoed21	shads15	stone95
13	passwoed	passw0ld	pazzword	neverman10	suckymon40	samerman10	shoes25	sheest5	stong05
14	passw00d	passw00d	pa55w0rd	hoterman10	supermon40	supercan10	stoed22	shars05	stoney5
15	passw0ld	passwoed	paztword	astarman10	sundoman40	supermin00	scoet24	shars95	stone25
16	passwuld	p@ssw0rd	pazsword	loverman20	survaman40	severman10	stoed23	shals15	smokey5
17	passw03d	passw03d	padsword	deverman10	summyman40	superman40	stoed20	shers05	stone15
18	passweed	passwuld	passwwrd	katerman10	supieman40	superman70	scoet29	shads25	stone05
19	passw0ld	passw0ld	pas2word	waterman10	sundymon40	saverian10	stoe123	shees35	sport15
20	passwitd	passweed	pardward	severman10	succemin40	superman50	stoe122	shars45	shoes25
Guessing success rate	1.87×10^{-2}	1.88×10^{-2}	1.89×10^{-2}	9.24×10^{-8}	2.93×10^{-6}	4.01×10^{-6}	3.08×10^{-7}	5.50×10^{-6}	9.24×10^{-6}

[†]Training on 10M Rockyou. C2C-impact=Character-to-character impact. A **bold** string means that the generated password hits (i.e., the same as) the target password.

TABLE XIX
BASIC INFORMATION ABOUT OUR PASSSEQ'S NETWORK STRUCTURE.

Layer Name	Description	Output shape
Encoder		
Embedding	Embeds the input sequences into dense vectors.	[batch-size, seq-length, embedding-size (128)]
Bi-LSTM×3	Processes the embedded input sequence.	[batch-size, seq-length, hidden-size (256)*2]
Linear	Transforms the encoder states to the decoder states.	[batch-size, seq-length, hidden-size (256)]
Decoder		
Embedding	Embeds the target input at each decoding step into a dense vector.	[batch-size, 1, embedding-size (128)]
LSTM cell	Processes the embedded target input and generates hidden states.	[batch-size, hidden-size (256)]
Linear (Attention)	Transforms the encoder output to calculate attention scores.	[batch-size, seq-length, hidden-size (256)]
Linear (Concatenation)	Combines attention context and decoder hidden state.	[batch-size, hidden-size (256)]
Linear (Output)	Transforms the decoder hidden state to the output vocabulary space.	[batch-size, vocab_size (95)]

and the much less frequent superman40 (rank 1,117,959) and shoes25 (rank 1,118,029). These examples illustrate how different masking strategies influence the template-guided decoding process. Building on these observations, Table XVII summarizes the overall success rates under identical evaluation settings, showing that within 5×10^8 guesses the overall ranking is C2C-impact>hotword-0.2>random-0.5. The C2C-impact method achieves the best results by modeling fine-grained dependencies between characters, while hotword masking benefits from preserving common human-chosen structural patterns. Random masking, despite its simplicity, can still perform competitively because it exposes PassSeq to a wide variety of template structures when trained on sufficiently large datasets. Taken together, the three strategies capture complementary forms of structural information (global template diversity in random masking, human composition patterns in hotword masking, and character-level dependency modeling in C2C-impact masking), which collectively strengthen PassSeq's ability to learn and generate plausible passwords.

D. Details of our maximum likelihood estimation

Let pw^* be a password that conforms to the template PT^* . Denote n as the number of unique passwords in the set Γ_{PT^*} that conform to PT^* , and m as the number of passwords in Γ_{PT^*} that also satisfy $p(\text{pw}) > p(\text{pw}^*)$. We want to estimate m . Let $p = m/n$, which is the probability that a uniformly sampled password from Γ_{PT^*} satisfies the condition $p(\text{pw}) > p(\text{pw}^*)$.

Define a sampling experiment $\mathcal{S}$: Uniformly sample one password pw from Γ_{PT^*} . If $p(\text{pw}) > p(\text{pw}^*)$, record a success. Repeat experiment $\mathcal{S}$ for N times. Let the number of successes be x , which follows a binomial distribution:

$$x \sim \text{Binomial}(N, p), \quad f(x|N, p) = \binom{N}{x} p^x (1-p)^{N-x}. \quad (24)$$

Define experiment $\mathcal{T}$ as repeating $\mathcal{S}$ N times. Repeat $\mathcal{T}$ for t times, and denote the sequence of success counts as $\{x_1, x_2, \dots, x_t\}$. The likelihood function is:

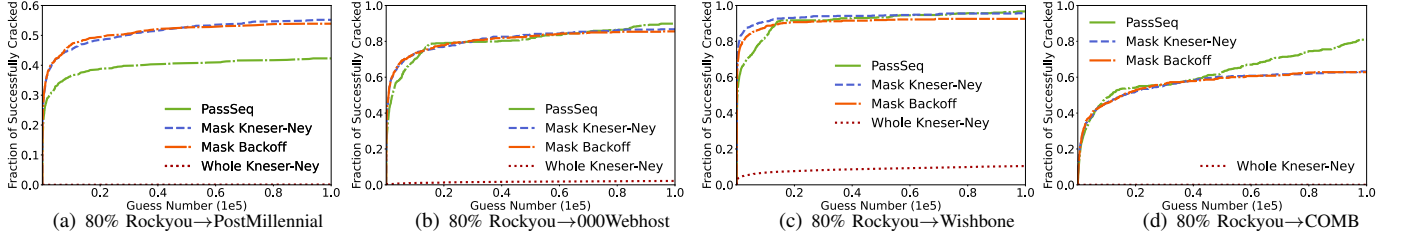


Fig. 14. Mask guessing with different character leakage probabilities (i.e., the type-2 attacker $\mathcal{A}_{\text{stat}}$). In these scenarios, the password is divided into three equal-length regions (front, middle, back) with per-character disclosure probabilities of 30%, 50%, and 70%, respectively, reflecting the monotonic increase in thermal residue observed in prior studies [4], [6]. When a victim's password characters are probabilistically leaked, the guessing success rate of $\mathcal{A}_{\text{stat}}$ can be 7.85-40.59 times higher than in scenarios without any character leakage (i.e., the results of the whole Kneser-Ney model).

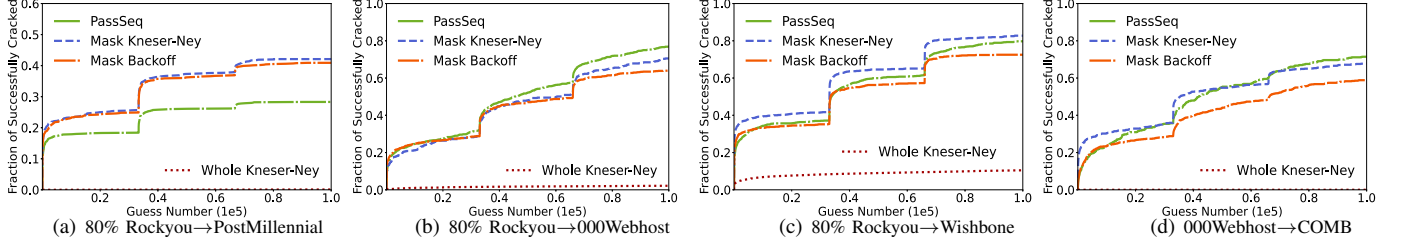


Fig. 15. Mask guessing without the password length information (i.e., the $\mathcal{A}_{\text{length}}$ attacker). When $\mathcal{A}_{\text{length}}$ is unaware of the victim's password length (but knows 50% characters with certainty), her success rate can increase by 5.93-34.59 times compared to whole password guessing (i.e., whole Kneser-Ney).

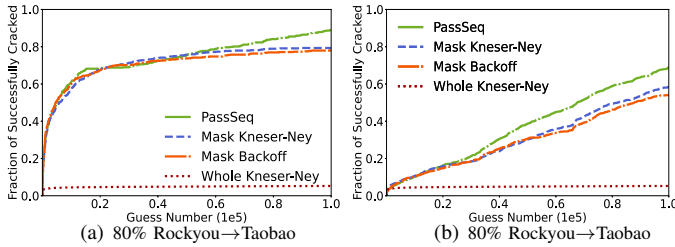


Fig. 16. Mask guessing under different character leakage probabilities (i.e., the type-2 attacker $\mathcal{A}_{\text{stat}}$) and without password length information (i.e., the $\mathcal{A}_{\text{length}}$ attacker), evaluated in cross-lingual scenarios where the training and testing datasets are from different languages. The experimental setup is consistent with Fig. 14 and Fig. 15, respectively.

$$\begin{aligned}
 L(p) &= \log \left(\prod_{i=1}^t \binom{N}{x_i} p^{x_i} (1-p)^{N-x_i} \right) \\
 &= \sum_{i=1}^t \log \binom{N}{x_i} + \sum_{i=1}^t x_i \log(p) + \sum_{i=1}^t (N-x_i) \log(1-p) \\
 &= \sum_{i=1}^t \log \binom{N}{x_i} + \left(\sum_{i=1}^t x_i \right) \log(p) + \left(tN - \sum_{i=1}^t x_i \right) \log(1-p).
 \end{aligned} \tag{25}$$

Taking the derivative of $L(p)$ with respect to p gives:

$$\frac{\partial L(p)}{\partial p} = \frac{\sum_{i=1}^t x_i}{p} - \frac{tN - \sum_{i=1}^t x_i}{1-p}. \tag{26}$$

Setting $\frac{\partial L(p)}{\partial p} = 0$ and solving yields the MLE:

$$\tilde{p} = \frac{1}{Nt} \sum_{i=1}^t x_i. \tag{27}$$

Unbiasedness: The expectation of $\tilde{p}$ is

$$E[\tilde{p}] = \frac{1}{Nt} \sum_{i=1}^t E[x_i] = \frac{tpN}{Nt} = p. \tag{28}$$

Variance: The second derivative of the log-likelihood is:

$$\frac{\partial^2 L(p)}{\partial p^2} = -\frac{\sum_{i=1}^t x_i}{p^2} - \frac{tN - \sum_{i=1}^t x_i}{(1-p)^2}. \tag{29}$$

The Fisher Information is the negative expectation:

$$\mathcal{I}(p) = -E \left[\frac{\partial^2 L(p)}{\partial p^2} \right] = \frac{tN}{p(1-p)}. \tag{30}$$

The Cramér-Rao bound (inverse Fisher Information) gives:

$$\text{Var}(\tilde{p}) = \frac{p(1-p)}{tN} \leq \frac{1}{4tN}. \tag{31}$$

Conclusion: The estimator $\tilde{p}$ is unbiased and converges to p in probability. The estimate of m is: $\tilde{m} = \tilde{p} \cdot n$.

E. Sampling from non-normalized password weight functions

In practice, some password models (especially those integrating multiple priors such as templates, side-channel information, or PII) do not explicitly provide a normalized probability distribution, but only assign each password a positive weight used for ranking. As our MLE-based guess number estimation requires only i.i.d. samples from the induced normalized distribution, we show that correct sampling remains possible via rejection sampling.

Non-normalized password weight function. Let PW denote the password space, and let $f : \text{PW} \rightarrow \mathbb{R}^{>0}$ be a positive weight function assigned by a password model. For a finite set $S \subseteq \text{PW}$, define

$$f(S) = \sum_{\text{pw} \in S} f(\text{pw}). \tag{32}$$

If $f(S) \neq 1$, we call f a *non-normalized password weight function* over S . Let

$$C = \max_{\text{pw} \in S} f(\text{pw}), \tag{33}$$

and assume that $C < \infty$. Our goal is to sample passwords from S according to the normalized distribution

$$\Pr(\text{pw}) = \frac{f(\text{pw})}{f(S)}. \tag{34}$$

