



**Politecnico
di Torino**

Politecnico di Torino

Quantum Engineering

A.a. 2025/2026

Graduation Session October 2025

Twin-Field Quantum Key Distribution on the telecommunication optical fiber network

**Analysis and experimental implementation of the asymmetric
SNS protocol**

Supervisors:

Prof. Scotognella
Dott.ssa Clivati

Candidate:

Francesco Magliano

Abstract

This thesis, developed during an experimental research project at the Italian National Metrology Research Institute (INRIM) in Turin, focuses on the practical implementation of the Twin-Field Quantum Key Distribution (TF-QKD) protocol. This approach is promising for realizing quantum-secure communication over long distances, as it suffers less from transmission loss compared to more conventional QKD protocols. However, it requires ultrastable lasers and fully coherent photon propagation over noisy channels.

The main objective of my research was to address the key technical challenges that have so far impeded its long-distance application in real-world optical fiber networks, with a particular focus on maintaining phase coherence.

After introducing the fundamentals of QKD, this work presents a theoretical analysis of the phase coherence problem, showing the impairments associated to the application of TF-KD in real-world optical networks. Based on the analyzed criticalities, we present an experimental architecture that applies the dual-band stabilization technique to an Independent laser scheme to detect and compensate for fiber noise in real-time without affecting quantum transmission. In particular, this approach uses commercial technologies, such as narrow-linewidth fiber lasers, to replace the complex and expensive ultra-stable lasers based on Fabry-Pérot cavities and vacuum systems used in previous demonstrations. This allowed for the realization of a compact and transportable setup suitable for the real-world application of TF-QKD in modern telecommunication networks.

During my experimental work, I had the opportunity to characterize the optical components of the setup. This included an analysis of the laser sources to evaluate Raman scattering as a function of their power, thereby limiting noise. I also characterized the phase and intensity modulators by evaluating their Extinction Ratio (ER), and the detectors in terms of their dark counts. For all components, their adequacy for the protocol's critical requirements was verified. Furthermore, specific control software was developed for the advanced Field-Programmable Gate Arrays (FPGAs), which are used within each communication node for intensity/phase modulation and temporal synchronization. This latter function is accomplished by also leveraging advanced protocols like White Rabbit for nanosecond-order synchronization

The thesis also includes an estimation of the expected secret key rate for a TF-QKD implementation on a deployed fiber network under realistic conditions, using the decoy-state method. This simulation is based on a complete mathematical analysis of the asymmetric SNS protocol variant, which is considered the most robust and

secure TF-QKD implementation to date.

Overall, the setup is distinguished by its robustness and compactness. By employing a complex optical system based on four active feedback controls, the setup demonstrates its ability to effectively compensate for fast and slow phase drift, polarization mismatch, and temporal delay between the twin fields. These features make it suitable for TF-QKD implementation on optical fiber links of up to 200 km in noisy environments, ensuring effective fast phase control for over 4 s with residual noise correction occurring in windows of about 100 ms.

Ultimately, this work validates the setup as a highly promising solution, paving the way for the system's integration into future global quantum networks.

Table of Contents

List of Figures	IV
1 Basic Notions of Classical Cryptography	1
1.1 Introduction	1
1.2 Symmetric Cryptography	3
1.2.1 Key Distribution	8
1.2.2 The Random Nature of the Key	10
1.2.3 Unconditional Security and the One-Time Pad	11
1.3 Principles of Asymmetric Cryptography	13
2 Quantum Cryptography	16
2.1 Quantum 2.0	16
2.1.1 The Impact of Quantum Computing on Modern Cryptography	17
2.1.2 The Advent of Quantum and Post-Quantum Cryptography .	18
2.2 Quantum Cryptography	20
2.2.1 Origin of Quantum Cryptography	20
2.2.2 Introduction to DV-QKD	21
2.2.3 Generic Structure of a Prepare-and-Measure QKD Protocol	27
2.2.4 Overview of DV-QKD protocols	32
2.2.5 Photon Number Splitting attack and Decoy states	37
3 TWIN-FIELD QKD PROTOCOL	41
3.1 Rate-Distance Limits of Current QKD Protocols	41
3.2 TF-QKD: A Fundamental Resource for Next-Generation Long-Distance QKD	42
3.2.1 TF-QKD Performance	43
3.3 Operating Scheme of TF-QKD	46
3.3.1 Basic Operating Principles	46
3.3.2 Evolution toward the TF-QKD setup	48
3.4 Properties of TF-QKD	50
3.4.1 TF-QKD vs MDI-QKD	50

3.4.2	Coordinated Twin-Field Phase Randomization	51
3.4.3	Derivation of TF-QKD Key Rate	53
3.5	Secure TF-QKD Protocols: SNS and CAL	54
3.5.1	QKD Protocol Security Definitions	54
3.5.2	Introduction to SNS and CAL Protocols	55
3.5.3	Asymmetric SNS Protocol Implementation	60
4	TF-QKD: Challenges and Real-World Implementations	73
4.1	The Problem of Phase Coherence	73
4.1.1	TF-QKD Critical Assumptions	73
4.1.2	Sources of Phase Fluctuation	74
4.1.3	Impact of Phase Stability on QBER	75
4.2	Real-World Phase Stabilization for TF-QKD	76
4.2.1	Spectral Phase Analysis	77
4.2.2	Common-Laser Scheme	77
4.2.3	Common Laser Scheme with Dual-Band Stabilization	82
4.2.4	Independent Laser Scheme	88
4.2.5	Independent Laser Scheme with Dual Band Stabilization . .	91
4.2.6	Comparative Analysis and Critical Considerations	93
5	Our Experiment	95
5.1	Introduction	95
5.2	Experimental Setup Overview	95
5.2.1	Architecture and Core Principles	95
5.2.2	Active Feedback and Control	97
5.3	Fast Phase Drift Cancellation System	100
5.3.1	Schematic of Alice and Bob's Enclosure	100
5.3.2	Charlie's node Basic Schematic	108
5.3.3	Analysis of Phase Stability	114
5.4	Slow Phase Drift Cancellation System	121
5.4.1	Experimental Setup and Analysis of Slow Phase Drift	122
5.4.2	Strategy for Slow Phase Drift Correction	124
5.4.3	Analysis of Slow Phase Drift Compensation	127
5.5	Temporal Realignment Procedure	129
5.5.1	Experimental Setup and Strategy	131
5.5.2	Automation of Realignment with Cross-Correlation	138
5.6	Building a QKD System: Components and Their Characterization .	141
5.6.1	Sources	142
5.6.2	Detectors	144
5.6.3	Raman Noise Analysis	147
5.6.4	Encoder and Signal Modulation	152

5.6.5	Variable Optical Attenuator	161
5.6.6	Characterization of the Single-Photon Regime for Phase Re- alignment	162
6	Code and Simulation for the Asymmetric SNS-TF-QKD Protocol	167
6.0.1	Physical Layer Simulation and Data Generation	168
6.0.2	Protocol Post-Processing and Security Analysis	173
7	Conclusions and Future Developments	179
A	Practical Procedure for Phase-Locking	182
B	Experimental parameters and technical details	185
C	Technical Steps for Delay Setting	187
D	Simulation Simplifying Hypotheses	190
E	Comparison of typical properties of APDs and SNSPD	192
	Bibliography	194

List of Figures

1.1	An overview of the field of Cryptology [1].The diagram illustrates Cryptology's main branches: <i>Cryptography</i> , the science of creating secure messages, and <i>Cryptanalysis</i> , the science of breaking cryptographic systems. Furthermore, it illustrates the main categories within cryptography, distinguishing between <i>symmetric</i> , <i>asymmetric ciphers</i> and also including <i>protocols</i> . It finally breaks down symmetric ciphers into their two main sub-categories: <i>block ciphers</i> and <i>stream ciphers</i>	2
1.2	The basic principle of symmetric-key encryption [1]. Alice encrypts her original message, the plaintext (x), into a ciphertext (y) using an encryption function (e) and a shared secret key (k). She sends the ciphertext to Bob via an insecure channel. Upon receiving the ciphertext, Bob uses the same secret key (k) and the inverse function e^{-1} . to decrypt the message and recover the original plaintext (x). .	3
1.3	A schematic of a <i>Symmetric-Key Cryptosystem</i> [1]. Alice encrypts the plaintext (x) into a ciphertext (y) using an encryption function (e) and a shared secret key (k). This key is transmitted to Bob via a separate secure channel. The ciphertext is sent over an insecure channel, where it can be intercepted by an eavesdropper, Eve. Upon receiving the ciphertext, Bob uses the same key (k) and a decryption function (d) to recover the original message.	4
1.4	Truth table of the XOR operation [1] between a bit of the plaintext x_i and a bit of the stream key s_i	5
1.5	A comparison of <i>Stream</i> and <i>Block ciphers</i> [1]. The diagram illustrates the fundamental difference between the two symmetric-key encryption methods: a Stream Cipher encrypts data one bit at a time, transforming the input bit stream $(x_0, x_1, \dots)$ into an output stream $(y_0, y_1, \dots)$; while a Block Cipher encrypts data by processing it in fixed-size blocks of b bits, transforming an input block $(x_0, \dots, x_b)$ into an output block $(y_0, \dots, y_b)$	6

1.6	The basic schema for a synchronous and asynchronous stream cipher[1]. The key stream generator, initialized with a secret key (k), produces a keystream (s_i). This is used with a bitwise XOR operation ($\oplus$) to encrypt the plaintext bit (x_i) into the ciphertext bit (y_i). The dotted line represents the feedback loop in an asynchronous stream cipher, where the keystream depends on a previous ciphertext bit, while in a synchronous stream cipher, it depends only on the initial secret key.	7
1.7	Basic scheme of a Block Cipher and the Data Encryption Standard (DES) algorithm [1]. The left side shows the general principle of a block cipher, which processes a fixed-size block of plaintext and a secret key to produce a block of ciphertext. The right side illustrates the specific parameters of the DES algorithm, which encrypts a 64-bit plaintext block using a 56-bit key to generate a 64-bit ciphertext block.	8
1.8	An illustration of the <i>Key Distribution Problem</i> in a Symmetric Cryptosystem. For a complete network graph with n nodes, the number of required secret key pairs is given by the formula: $C_2 = \frac{n(n-1)}{2}$. [2] This figure shows an example with five nodes, requiring a total of 10 unique key pairs. This rapid increase in the number of keys to be managed and stored securely represents a significant challenge for symmetric cryptography, making it impractical for large-scale networks.	9
1.9	The basic schema for Public-Key (asymmetric) encryption[1]. The recipient, Bob, has a key pair composed of a public key (k_{pub}) and a private key (k_{pr}). Alice encrypts her plaintext (x) using Bob's public key, creating the ciphertext (y). Bob then decrypts the message using his unique private key. The public key can be openly shared, while the private key remains secret, eliminating the need for a secure key distribution channel.	13
2.1	Representation of the Bloch Sphere [12]	21
2.2	Different degrees of freedom for encoding quantum information using single photons [15]. The figure illustrates how a qubit can be represented by various physical properties of a photon, including its Fock state (number of photons), Time-Bin (arrival time), Polarization (horizontal or vertical), Dual-Rail (path), and Frequency. These properties are manipulated to encode the logical states $ 0\rangle$ and $ 1\rangle$ for quantum communication.	24

2.3	Basic schema of Quantum Key Distribution (QKD)[2]. Alice prepares quantum states and sends them to Bob via a quantum channel. Bob performs quantum state detection. An eavesdropper, Eve, may attempt to intercept the quantum channel. Crucially, an authenticated classical channel is also used between Alice and Bob for tasks like basis reconciliation and error correction, ensuring the security of the shared key even in the presence of Eve.	27
2.4	Basic schema of the Measurement-Device-Independent Quantum Key Distribution (MDI-QKD) protocol. Alice and Bob each prepare quantum states and send them to a central, untrusted node, Charlie. Charlie performs a Bell State Measurement (BSM) on the incoming signals and publicly announces the result. This protocol's security is guaranteed because the BSM result reveals a correlation between Alice's and Bob's bits without Charlie learning the actual values, making the protocol robust against <i>Detector Side-Channel attacks</i> . .	36
3.1	A homogeneous comparison between theoretical limits (lines) and experimental results (symbols) for different QKD schemes, reproduced from Lucamarini et al. [29]. The graph is normalized to a standard optical fiber with an attenuation coefficient of $\alpha = 0.2 \text{ dB km}^{-1}$. The theoretical limits shown are: I, MDI-QKD with decoy states [27]; II, QKD with decoy states [28]; III, single-photon QKD [8]; and IV, The maximum theoretical secret key capacity (SKC) is indicated by the PLOB bound (dashed yellow line) [8]; ideal TF-QKD [29], and the single quantum repeater limit [8]. The experimental results, indicated by squares, triangles, and circles, correspond to discrete variable QKD, CV-QKD, and MDI-QKD, respectively, numbered in corresponding chronological order. A solid black line represents the key rate from a real TF-QKD experiment, which, along with the ideal TF-QKD key rate (dashed black line), surpasses the repeater-less limits in the shaded area. This demonstrates that the TF-QKD key rate scales as the square root of channel transmittance (η), which is similar to a single quantum repeater, offering a significant advantage over other QKD protocols. The inset table shows the key simulation parameters for the ideal bounds and the TF-QKD experiment: P_{dc} represents the detector dark count rate, η_{det} is the detector efficiency, e_{opt} is the optical misalignment error, and f is the error correction efficiency factor.	44

3.2	A basic operating scheme of the TF-QKD protocol [29]. Alice and Bob use distinct laser sources (LS) and modulators to prepare weak optical pulses with encoded phase information. They send these pulses through two independent optical fibers of length L to an untrusted central node, Charlie. At Charlie's station, the two pulses interfere on a beam splitter (BS), and the outcome is measured by two single-photon detectors (D0 and D1), which are then publicly announced.	46
3.3	A detailed operating scheme of the Twin-Field Quantum Key Distribution (TF-QKD) protocol [29]. Alice and Bob each use a laser source (LS), intensity modulator (IM), and phase modulator (PM) to generate and encode weak coherent states with intensity (μ_a, μ_b) and phase information (φ_a, φ_b) . These pulses are sent through separate optical fibers to an untrusted node, Charlie. At Charlie's station, the pulses interfere on a beam splitter (BS), and the outcome is measured by two detectors (D0 and D1). An additional phase modulator (PM) at Charlie's station is used for phase compensation, while the random number generators (RNG) at Alice and Bob's sites are used to choose phases and bases. The figure on the right illustrates the different <i>phase slices</i> , which correspond to distinct phase values on the Bloch sphere and they are used during the <i>Reconciliation Twin-Field Phase</i> step.	47
3.4	The evolution toward the final TF-QKD setup[29]. (a) A typical phase-based QKD protocol where a single source (LS) sends two pulses through an asymmetric Mach-Zehnder interferometer to Bob via a single quantum channel of length L . (b) An intermediate scheme where the two pulses travel on separate channels of the same length to Bob, conceptually leading to the TF-QKD idea of twin fields. (c) The final TF-QKD scheme where two separate sources from Alice and Bob send pulses over distinct channels to a central, untrusted node (Charlie). This setup doubles the effective distance between Alice and Bob to $L_{\text{tot}} = 2L$ for the same count rate, allowing for longer-distance quantum key exchange.	48
3.5	Simulated key rates for BB84, SNS-AOPP and CAL protocols as a function of channel loss, using a setup with ultra-stable independent lasers and SPAD detectors. This result is reported in ref.[16].	59

- 3.6 Schematic diagram of the setup for the SNS protocol. This diagram shows the setup for an SNS (*Send-No-Send*) protocol, where we indicate IM as intensity modulator; PM as phase modulator; BS as beam splitter; D1 & D2 as single-photon detectors in Charlie's measurement station. Unlike other protocols, in this diagram it is possible to implement an asymmetric version of the SNS, in which Alice and Bob are not constrained to use the same source parameters. They can therefore independently choose the intensity of the signal states (s_A, s_B) and the probability of sending coherent pulses in the Z windows (ϵ_A, ϵ_B), without them having to be equal. 61
- 3.7 This graph shows the optimized key rates in relation to the transmission distance between Alice and Charlie obtained from different simulations reported in ref. [36]. The results of three different SNS protocols are compared: *asymmetric SNS*, *modified SNS*, and *original SNS*, while keeping the length difference between Alice and Bob's channels fixed at 50 km. It is important to note that these results consider the effect of the finite size of the exchanged data. As shown, the performance of the asymmetric SNS protocol is significantly better than that of the original SNS, especially when the length difference between the channels is large. Furthermore, it can be observed that the asymmetric SNS protocol surpasses the linear theoretical limits for repeaterless QKD, such as the famous PLOB bound, demonstrating the significant advantage of the TF-QKD protocol over its original MDI-QKD version. 62
- 4.1 A schematic of the experimental setup used in [29] to measure phase drift in a Mach-Zehnder interferometer with long optical fibers. A continuous-wave laser is split, and the two pulses travel along separate paths to an interfering beam splitter. The setup includes polarization controls, variable optical attenuators, and variable-length optical fibers to simulate different conditions. 73

4.2	Ideal Scheme of the TF-QKD Protocol [16]. Alice and Bob generate quantum states (QS) using two local lasers, attenuated to a single-photon level via variable optical attenuators (VOA). They modulate their respective quantum signals using intensity and phase modulators based on the specific TF-QKD protocol variant chosen (Section 3.3). In this ideal scenario, Alice's and Bob's quantum lasers have perfectly equal frequencies ($\nu_A = \nu_B$) and maintain total phase coherence. The signals are sent via two quantum channels of equal length (L) toward Charlie's central station, where they are measured by single-photon detectors (D_0 and D_1). This theoretical scheme automatically satisfies the two critical real-world TF-QKD requirements (Section 4.1.1).	76
4.3	Common-Laser Scheme for the TF-QKD protocol [16]. In this scheme, a reference laser (<i>Ref. laser</i>) positioned at the central station of Charlie provides a stable frequency (ν_R) to Alice and Bob through a dedicated <i>service channel</i> . This ensures that Alice's and Bob's local lasers have identical frequencies ($\nu_A = \nu_B = \nu_R$), satisfying the first critical assumption of TF-QKD. Furthermore, to address dynamic phase noise introduced by environmental variations (e.g., temperature, vibrations), the protocol employs an interleaving mechanism [43]. This periodically interrupts QKD (QS pulses) to switch to a classical regime where Charlie sends more intense reference pulses (Ref. pulses). By measuring the phase shift of these pulses, Charlie can accurately determine the relative phase error between Alice's and Bob's twin fields and apply the necessary real-time correction.	78
4.4	Time Scheme for Classical-Quantum Interleaving in TF-QKD [12]. Due to the rapid increase in relative phase noise and the consequent degradation of the QBER, the unstabilized TF-QKD protocol cannot operate continuously. Communication is divided into cycles: short quantum windows (q.s., quantum state), typically limited to $\sim 50 \mu\text{s}$ or less [43], are dedicated to sending the photonic signals for key generation. These windows are interleaved with classical reference periods (ref., reference), during which more intense (classical) light pulses are sent for monitoring and active realignment of the relative phase $\Delta\phi$ between Alice and Bob. This need for frequent realignment drastically limits the useful time for key extraction, reducing the overall efficiency of the SKR.	80

4.5	The Common Laser scheme with dual-band stabilization for TF-QKD. [44] A reference laser (ν_R) and a sensing laser (ν_S) are used to stabilize the phases of Alice's (ν_A) and Bob's (ν_B) local lasers, which are used for quantum state (QS) preparation. The setup employs two distinct fiber types: <i>QKD fibers</i> for quantum signal transmission and <i>service fibers</i> for laser synchronization. Charlie's station includes a noise detection and cancellation (NDC) system, a photodiode (PD) for phase monitoring, and an actuator (act.) for active fiber stabilization, ensuring high secret key rates over long distances.	82
4.6	A real-world TF-QKD common laser scheme based on <i>Dual-Band stabilization</i> , implemented by INRiM [41]. This setup uses a reference laser (ν_R) and a sensing laser (ν_S) to stabilize the local QKD lasers. The sensing laser travels a round trip, first on a service fiber and then on a QKD fiber, accumulating phase information that is used in Charlie's station to actively correct for fiber noise via an Acousto-Optic Modulator (AOM).	85
4.7	A schematic of the scientific experiment for TF-QKD using real optical fibers from the Italian Quantum Backbone [41]. Alice and Bob's terminals are located in Bardonecchia and Santhià, and Charlie's terminal is in Turin. This setup demonstrates the feasibility of dual-band stabilization for long-distance QKD over approximately 200 km, achieving a QBER of less than 1%.	87
4.8	First proposed scheme based on the independent-lasers approach to TF-QKD, using ultrastable laser sources (u.s. lasers)[16]. This scheme shows the setup with Alice, Bob, and the central Charlie node, including their respective photon sources and detectors. . . .	88
4.9	Scheme of the independent-lasers approach with dual-band stabilization for TF-QKD [16]. This method utilizes two ultrastable laser sources (ν_A, ν_B) and two sensing lasers ($f_{0A} = f_{0B} = f_0$) to enable high-bandwidth active phase noise cancellation on the optical channels from Alice and Bob to Charlie.	91

4.10	Real experimental setup used in [39] as an example of Independent Lasers with Dual-Band Stabilization. Alice and Bob, each equipped with an ultrastable laser (ν_A, ν_B), utilize a Phase Modulator (PM) to locally generate an optical frequency comb. Two spectral lines are extracted: the Quantum Signal (λ_q , encoded via the Encoder) and the Sensing Signal (λ_s , or sideband f_0). These signals travel through independent fibers to Charlie. At Charlie, the interference of the Sensing Signals ($\lambda_{sA,B}$) is detected by D_c to generate a real-time feedback signal. This signal drives a PID controller, which uses a PM on Alice's path to correct the phase drift of the reference channel, thereby automatically stabilizing the phase of the quantum channel ($\lambda_{qA,B}$) due to their fixed phase relationship. The corrected Quantum Signals then interfere, producing clicks on the Single-Photon Detectors (D_0, D_1).	92
5.1	Diagram of our experimental setup, which addresses key challenges of long-distance TF-QKD in realistic fiber infrastructures. The architecture is based on an independent laser scheme with a dual-band stabilization strategy (Solution 2.1). It features a pair of lasers at each terminal (Alice and Bob): a sensing laser ($\lambda_s = 1542.1$ nm) for phase compensation and a quantum laser ($\lambda_q = 1543.3$ nm) for key exchange. These lasers are mutually phase-locked at each terminal using a PLL. The signals are then sent to a central Charlie node, which uses a third PLL to correct fast phase drift. The system's effectiveness is further enhanced by four active control loops for fast phase drift, slow phase drift, time realignment, and polarization mismatch compensation, ensuring a robust and secure protocol.	96
5.2	Portion of the experimental setup utilized for our TF-QKD experiment conducted at the INRIM laboratories in Turin. The three network nodes—Alice, Bob, and Charlie—were collocated in separate enclosures within a laboratory where the ambient temperature was actively stabilized at 23 °C. The terminal nodes were connected to the central node via dedicated optical fibers.	97
5.3	Schematic of the experimental setup for Alice's (or Bob's) enclosure, illustrating the local phase-locking mechanism.	101

5.4	This graph illustrates how the energy of a phase-modulated signal is redistributed. As the modulation index (β) increases, the amplitude of the carrier frequency (J_0) decreases, while energy is transferred to higher-order harmonics ($J_1, J_2, \dots$). This physical principle means that with a higher modulation index, more discrete 'tines' of the frequency comb will appear and become more prominent.	102
5.5	This figure shows the optical frequency comb generated by the NKT sensing laser after a phase modulator (PM). The comb extends over 140 GHz, with sidebands labeled from $m = 0$ to $m = 4$ that are equally spaced by the 35 GHz frequency of the modulating RF signal. The red line represents the quantum laser.	105
5.6	Spectrum of the beat signal between the quantum laser and the $m = 4$ sideband of the sensing laser, maintained stably at 10 MHz by the PLL.	106
5.7	Encoder box: IM _{1,2} = two intensity modulators in cascade; PM = phase modulator; VOA = variable optical attenuator.	107
5.8	The diagram shows the basic schema of the central Charlie node's enclosure for a Twin-Field Quantum Key Distribution protocol. Its primary role is to measure and correct the total phase drift between the laser signals from Alice and Bob. The signals, with phases φ_A and φ_B , are combined, and a photodiode (PD) detects the beat signal. After being processed by a mixer and a low-pass filter, the output signal becomes proportional to $\sin(\varphi_A - \varphi_B)$. This output, which approximates the phase difference ($\varphi_A - \varphi_B$) for small fluctuations, is used to actively compensate for noise and maintain phase coherence, a critical requirement for the success of the QKD protocol.	109
5.9	Diagram of the experimental setup utilized to evaluate the effectiveness of the phase stabilization circuit.	115

5.10	<i>Combined Phase Power Spectral Density</i> (PSD) of the stabilized phase. This graph was generated by combining two distinct PSD datasets, each calculated using the <i>Welch method</i> with specific parameters: a larger number of data segments for the high-frequency analysis ($f > 1kHz$) and a smaller number of segments for low-frequency analysis ($f < 1kHz$) to accurately capture the noise. The result is a single curve showing the stabilized phase noise distribution over a broad frequency range, from approximately 1 Hz to 1 MHz. The clear suppression of high-frequency noise confirms the stabilization circuit's effectiveness in correcting fast phase drift. However, the residual low-frequency noise remains visible, which is likely attributable to the imperfect correlation between the noise experienced by the sensing and quantum signals due to the distinct fiber paths between them (see Section 5.4).	118
5.11	Cumulative Phase Standard Deviation (σ_ϕ). The graph shows the standard deviation of the phase between the two QKD lasers, calculated as a function of frequency. This result is obtained by numerically integrating the Phase Power Spectral Density (PSD), starting from high frequencies. The resulting curve represents the total contribution of phase noise across different timescales. The stabilization at approximately 0.12rad demonstrates the circuit's effectiveness in compensating for rapid phase fluctuations, ensuring the long-term stability required for the QKD protocol.	119
5.12	Experimental setup for evaluating the photon count rate over time. In this configuration, the interference of the quantum lasers from Alice and Bob is produced in the Charlie node, resulting in a click on the Single Photon Avalanche Diode (SPAD). The <i>CountRate Class</i> is used to process these clicks, producing a graph that allows for the precise evaluation of the effect of slow phase drift on the system.	122
5.13	Photon count rate per 100 ms time bin, illustrating the effect of residual phase drift on quantum signal interference. In this scenario, the count rate is highly fluctuating (non-deterministic), deviating from the predictable maximum/minimum rates expected in an ideal deterministic condition. This uncontrolled slow phase drift makes the reliable implementation of the TF-QKD protocol impossible, underling the necessity of addressing this residual noise for secure key exchange.	123
5.14	The evolution of the count rate for a temporal bin of 10 ms.	128
5.15	The evolution of the count rate for a temporal bin of 100 ms.	128
5.16	Code snippet used to add Coarse Delay to all modulators of Alice and Bob.	130

5.17	Representation of the experimental setup for temporal realignment. The core of this strategy is the use of the <i>Time Tagger Class</i> to build a histogram that visually represents the arrival times of Alice and Bob's pulses. By analyzing the position of the two distinct peaks on this histogram, Charlie's FPGA can calculate the relative delay between the two nodes. The FPGA then uses this value as feedback to apply precise coarse and fine delays, ensuring the temporal alignment necessary for the QKD protocol.	132
5.18	The Histogram class and its application for temporal synchronization in the experimental setup [50]. The down image illustrates how the class measures the time difference between a common <i>start signal</i> and multiple subsequent <i>click events</i> . In our experiment, the <i>start signal</i> simulates a White Rabbit's PPS, while the clicks represent the <i>quantum pulses</i> detected by the SPADs. The up image shows the resulting histogram, which accumulates these time differences into bins. The two distinct peaks visible in the histogram correspond to the arrival times of the pulses from Alice and Bob, allowing us to precisely measure and correct the mutual delay between the two nodes using Charlie's FPGA [50].	133
5.19	Histogram showing detected pulse counts before DC point optimization. The broader, lower peak (around 340 ns) corresponds to the optical pulses sent by Alice. This indicates a non-optimal DC bias setting for the intensity modulator, resulting in reduced signal transmission and a less clearly defined pulse at Charlie's detector. .	135
5.20	Histogram showing detected pulse counts after DC point optimization. The sharp, high peak (around 340 ns) demonstrates the successful optimization of the intensity modulator's DC bias at the Alice node. The temporal delay remains the same due to the fixed fiber length between Alice and Charlie, but the improved DC point results in a clearly defined, high-transmission pulse. This is crucial for subsequent temporal alignment and confirms that the system is operating at its point of maximum efficiency.	136

5.21	An example of a histogram for temporal realignment, showing two distinct peaks that represent the arrival times of pulses from the Alice and Bob nodes. The histogram was acquired over approximately 1 s, with a bin width of 100 ps. The separation between the peaks indicates the temporal delay that needs to be compensated using the FPGA in Charlie, which applies both <i>coarse</i> and <i>fine delay</i> adjustments. In this specific case, the measured delay between Alice and Bob is approximately 8 ns, which, considering the speed of light in a 1550 nm single-mode (SM) fiber (2×10^8 m/s), corresponds to a difference in fiber lengths (L_{A-C} and L_{B-C}) of approximately 1.6 meters. The width of each peak is approximately 2 ns, consistent with the expected duration of the modulated quantum pulses and, furthermore, it also confirms the correct temporal alignment of the cascaded intensity modulators (IM1,2) at their respective nodes, which are set with a fine delay. In this process, the characteristics of the detector, a SPAD, must be considered: in our case, the detector has an efficiency of 10% and a dead time of 10 ms after each click; however, potential Raman noise was not optimized in this particular measurement.	137
5.22	Illustration of <i>Cross-Correlation</i> for signal analysis. The four subplots demonstrate the correlation between two time signals, F and G. The top-left plot shows situation in which the signals are in-phase, while the top-right plot shows the corresponding cross-correlation result, with a red dot indicating the peak correlation. Similarly, the bottom-left plot shows the situation in which signals are in anti-phase, and the bottom-right plot shows the corresponding minimum correlation. In our specific application, we implement <i>Auto-Correlation</i> (where $F=G$) on the histogram of Alice's and Bob's pulse arrival times. This technique aims to automatically identify the relative temporal delay between the two quantum nodes by locating a distinct peak in the autocorrelation function, excluding the central peak at zero delay.	139
5.23	Autocorrelation result from a non-optimized, high-noise scenario. The high background noise obscures the smaller peaks, making it impossible to automatically identify the temporal delay.	140
5.24	Autocorrelation result from an optimized, low-noise scenario. The two distinct side peaks at ± 8 ns clearly indicate the precise temporal delay between the pulses from Alice and Bob.	140
5.25	Result of the temporal alignment. This histogram shows the final, single peak after successfully compensating for the relative delay between Alice's and Bob's signals.	141

5.26	The two types of lasers used in our <i>setup</i> : a narrow-linewidth fiber laser manufactured by NKT as a sensing laser, and a diode laser manufactured by RIO as a quantum laser.	142
5.27	The power spectra as a function of wavelength are shown for the NKT (1542 nm) and RIO (1543.3 nm) lasers.	143
5.28	Comparison of the phase noise spectra for the RIO and NKT lasers. This figure highlights how the phase noise of the NKT laser is significantly lower than that of the RIO laser, confirming its suitability as a narrow-linewidth source for the TF-QKD protocol.	145
5.29	Setup and operating parameters of the ID220 SPAD detector. The graph shows the user interface of the single-photon detector (SPAD), highlighting the configuration parameters used in our setup: a detection efficiency of 10% and a dead time of 10 μ s, essential values for system characterization and power budget.	146
5.30	A trace of the detector's background count rate over time, measured in counts per second, is shown in the graph. The graph illustrates the intrinsic background noise and dark photons, showing an average value of approximately 20 counts with an integration window of 100 ms and a standard deviation of about ± 10 counts. This measurement was performed using the implemented <code>Count_Counter</code> class.	147
5.31	Experimental data fitting with the Raman model. The graph shows the measured Raman power (blue points) as a function of fiber length for a fixed sensing laser launch power ($I_0 = 13$ dBm). A non-linear fitting algorithm was used to apply the theoretical model, $R(L) = \beta \cdot I_0 \cdot L \cdot e^{-\alpha L}$, to the experimental data. The best-fit curve (solid blue line) and the corresponding parameters β and α are displayed. The red circle indicates the maximum <i>Effective Length</i> at which the Raman noise power reaches its highest value before beginning to decrease due to fiber attenuation.	148
5.32	Raman noise as a function of fiber length for three different sensing laser launch powers: 13 dBm, 3 dBm, and -7 dBm. Each curve illustrates how the Raman noise scales with distance for a specific launch power. The red circles indicate the lengths at which the signal-to-noise ratio (SNR) of the beat note drops below a critical threshold of 30 dB. The graph demonstrates that even at low launch powers, an acceptable SNR can be maintained for long distances, suggesting that a total communication range of up to 200 km is feasible for the TF-QKD protocol in this setup.	149

5.33	A simplified representation of the encoder's modulation pattern. The signal features a 2 ns pulse followed by a 2 ns empty buffer, resulting in a total period of 4 ns.	153
5.34	Standard Modulator Characteristic [52]: A typical modulator operation where V_o is set at the point of maximum slope for linear modulation. A DC bias point is applied, and an RF signal modulates around it.	155
5.35	Our Modulator Characteristic with Inverted Logic: Our specific setup is calibrated such that the optimal DC point corresponds to minimum optical output. A negative voltage from the FPGA generates a high optical pulse (logical "1"), while 0V maintains minimum output (logical "0").	155
5.36	Modulation patterns generated by the twin FPGA boards for the intensity modulators (IM1 and IM2) and the phase modulator (PM). The identical intensity patterns for IM1 and IM2 (red and green lines) were implemented for a simulation of the asymmetric 4-decoy SNS TF-QKD protocol, displaying four distinct intensity levels: signal, vacuum=decoy0, decoy1 and decoy2 (Section 3.5.3). The phase modulator (PM) pattern (blue line) shows an example of a subset of the 16 distinct phase levels used for coordinated phase randomization (Section 3.2)	157
5.37	Characterization of the Variable Optical Attenuator (VOA) used in Alice and Bob's enclosures.	161
5.38	The graph shows the progressive increase of the detector count (orange line) as the voltage supplied to the analog VOA increases (corresponding to a decrease in optical attenuation). The count stabilizes at 5000 clicks per second (blue line), a value that confirms the achievement of the single-photon regime and coincides with the theoretical number of pulses emitted in a 100 ms interval.	164
5.39	Anomalous behavior of the SPAD detector in response to a high-intensity signal. The graph shows a count that exceeds the theoretical limit of 5000 clicks, indicating a non-ideal behavior of the quenching circuit and the subsequent inability of the detector to correctly discriminate multiple pulses.	165
A.1	Spectrum of the beat signal between the <i>sensing</i> lasers of Alice and Bob at 90 MHz.	183
B.1	Software interface for controlling the sensing lasers (NKT) of Alice and Bob	185

B.2	Software interface for controlling the quantum laser (RIO). The top panel shows the main parameters for current and temperature control, while the bottom panel shows the PID controller settings and the real-time feedback loop status.	186
-----	---	-----

Chapter 1

Basic Notions of Classical Cryptography

1.1 Introduction

As is often the case in the field of quantum technology, the first step towards a proper overview of a problem requires a brief look back at the world we define as "classical." Only then can we delve into how a technology naturally evolves from classical to quantum. To understand what we mean by quantum cryptography, we must start with the definition of classical cryptography. In its most general sense, cryptography can be defined as a mathematical procedure used to protect a message from unauthorized access. Cryptography can be traced back to 2000 BC, with the use of "non-standard" hieroglyphs by the Egyptian elite and was widely employed in ancient Greece and the Roman Empire, with the most famous examples being the Spartan scytale and the Caesar cipher.

Beyond these historical curiosities, we will focus on modern cryptography, attempting to provide a practical overview of the field. As shown in Fig.[1.1] , the most accurate and comprehensive term is *Cryptology*, which can be divided into two main sub-fields [1]:

- *Cryptography*: The science of creating a secret message. Its objective is to encrypt a message to make it secret and thus hide it from any unauthorized user attempting to intercept it.
- *Cryptoanalysis*: Conversely, this is the science of deciphering cryptographic systems to recover the original form of an encrypted message. It is also used to determine whether a cryptographic system is secure.

In this work, we will focus primarily on the cryptography branch, which, as

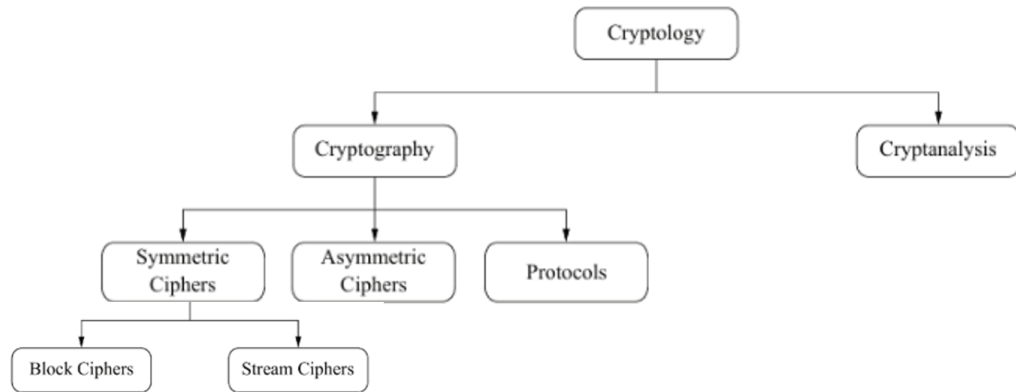


Figure 1.1: An overview of the field of Cryptology [1]. The diagram illustrates Cryptology's main branches: *Cryptography*, the science of creating secure messages, and *Cryptanalysis*, the science of breaking cryptographic systems. Furthermore, it illustrates the main categories within cryptography, distinguishing between *symmetric*, *asymmetric ciphers* and also including *protocols*. It finally breaks down symmetric ciphers into their two main sub-categories: *block ciphers* and *stream ciphers*

shown in the Fig.[1.1] , is further divided into three sub-fields [1]:

- *Symmetric Algorithms (or Private Key)*: These algorithms are based on the assumption that the two parties involved in a secret communication share a single secret, or private, key that is known only to them. This key is used for both encrypting and decrypting the message. It is noteworthy that all cryptography until the mid-20th century was based on these algorithms, which are still widely used today in many systems for data encryption and message integrity control.
- *Asymmetric Algorithms (or Public Key)*: In these algorithms, the two parties do not share the same secret key. Instead, one user possesses a secret key known only to them, while the other user has a public key that is not secret and can therefore be known to a potential attacker. These algorithms emerged in 1976 with key figures in modern cryptography such as Whitfield Diffie, Martin Hellman, and Ralph Merkle. Today, they are employed in numerous applications, including digital signatures, data encryption, and the distribution of a secret key for subsequent symmetric encryption.
- *Cryptographic Protocols*: While symmetric and asymmetric algorithms are the building blocks of modern cryptography, cryptographic protocols are the

fundamental elements that apply these algorithms to create various applications. A particularly important example today is the Transport Layer Security (TLS) protocol, the most widely used protocol in web browsers for secure communication over the internet.

As we will see in the following subsections, most modern cryptographic applications exploit a combination of symmetric and asymmetric algorithms, forming what are known as *hybrid schemes*. This approach aims to combine the strengths of both families of algorithms, thereby mitigating their respective weaknesses.

1.2 Symmetric Cryptography

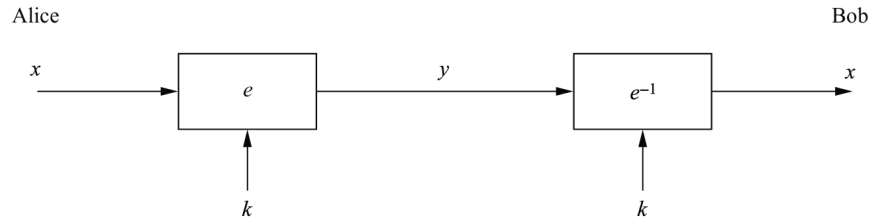


Figure 1.2: The basic principle of symmetric-key encryption [1]. Alice encrypts her original message, the plaintext (x), into a ciphertext (y) using an encryption function (e) and a shared secret key (k). She sends the ciphertext to Bob via an insecure channel. Upon receiving the ciphertext, Bob uses the same secret key (k) and the inverse function e^{-1} to decrypt the message and recover the original plaintext (x).

Introduction and definitions

This subsection aims to provide a more detailed introduction to the symmetric cipher scheme and to familiarize ourselves with the technical terms used in modern classical cryptography. As we will see, this type of scheme is what we will typically use for quantum cryptography, which will be introduced in Chapter 2.

As previously mentioned, symmetric cryptography, also known as secret-key or single-key cryptography, is characterized by a single secret key used by the parties for both encryption and decryption. In particular, Fig.[1.2] shows the typical scheme used to explain this concept, which will appear frequently in this work, involving two users, Alice and Bob, who want to communicate a secret message over an insecure channel, where there is a non-zero probability that a malicious

third party, Eve, could intercept the message.

The problem to be solved is simple: how can Alice and Bob exchange a message they want to keep secret, despite the possibility of an attacker intercepting and reading it?

It is useful at this point to introduce the technical terms typically used in the context of symmetric cryptography are as follows [1]:

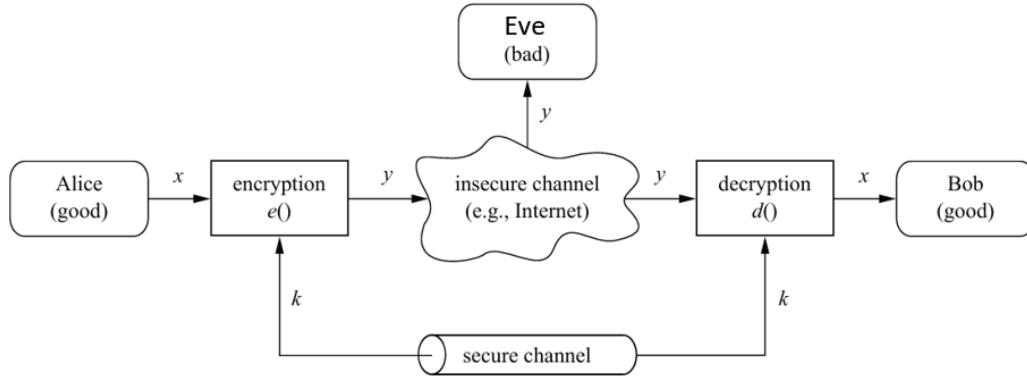


Figure 1.3: A schematic of a *Symmetric-Key Cryptosystem*[1]. Alice encrypts the plaintext (x) into a ciphertext (y) using an encryption function (e) and a shared secret key (k). This key is transmitted to Bob via a separate secure channel. The ciphertext is sent over an insecure channel, where it can be intercepted by an eavesdropper, Eve. Upon receiving the ciphertext, Bob uses the same key (k) and a decryption function (d) to recover the original message.

- $x = \text{Cleartext}$ or plaintext. This is the original message that Alice and Bob want to share and keep secret from Eve.
- $y = \text{Ciphertext}$. This is the encrypted message that Alice uses to protect it from an interception attack by Eve, known as *eavesdropping*.
- $k = \text{Secret key}$. This is the common secret key that Alice and Bob have shared before communicating the cleartext using a secure channel and with which they can protect the message during communication through the two phases of encryption and decryption (see Fig.[1.3]).

The following steps, illustrated in Fig.[1.3], show how symmetric cryptography effectively solves the communication problem introduced above:

1. Alice takes her original message, the cleartext (x), and encrypts it using a symmetric algorithm and a secret key (k).
2. The result is a secret message called the ciphertext (y).
3. Alice sends the ciphertext to Bob via the insecure channel.
4. Bob receives the ciphertext and decrypts it using the same secret key (k) that Alice used. This process reverses the encryption, restoring the original message (x).

The result of these operations is that, to an attacker like Eve, the ciphertext will appear as a random sequence of data, unrelated to the original message.

x_i	s_i	$y_i \equiv x_i + s_i \bmod 2$
0	0	0
0	1	1
1	0	1
1	1	0

Figure 1.4: Truth table of the XOR operation [1] between a bit of the plaintext x_i and a bit of the stream key s_i

Modulo 2 Addition and Its Role in Encryption The elementary XOR operation is the most widely used in modern cryptography. This comes from a fundamental property of this operation: perfect balance.

By observing the truth table of the XOR operation (Fig.[1.4]) and assuming for simplicity a simple scheme where each cleartext bit (x_i) is encrypted by performing the XOR operation with the corresponding bit of a stream key of the same length ($x_i \oplus s_i$), one can easily note that:

- If $x_i = 0$, the ciphertext bit after encryption (y_i) is equal to the key bit (s_i). If s_i is 0, the result is 0. If s_i is 1, the result is 1.
- If $x_i = 1$, the result (y_i) is the inverse of the key bit (s_i). If s_i is 0, the result is 1. If s_i is 1, the result is 0.

From this, we deduce a fundamental property: if the key is completely random, meaning each bit has a 50% probability of being 0 or 1, then through the XOR operation, the ciphertext bit (y_i) will also have a 50% probability of being 0 or 1. This characteristic ensures that our cryptographic system transfers the random nature of the chosen key to the ciphertext, as an attacker who intercepts and observes the ciphertext is unable to obtain any useful information about the original text. If they see a "1" or "0" bit in the ciphertext, it can correspond with equal probability to the encryption of a "0" or a "1" in the cleartext.

Furthermore, another important feature of XOR is its invertibility. This means that when we want to decrypt an encrypted message (y_i) with XOR, we just need to reapply the same operation with the same key (s_i) (Fig.[1.2])

At the single i -th bit level, we can write:

$$y_i \oplus s_i = (x_i \oplus s_i) \oplus s_i = x_i.$$

Stream Ciphers vs. Block Ciphers

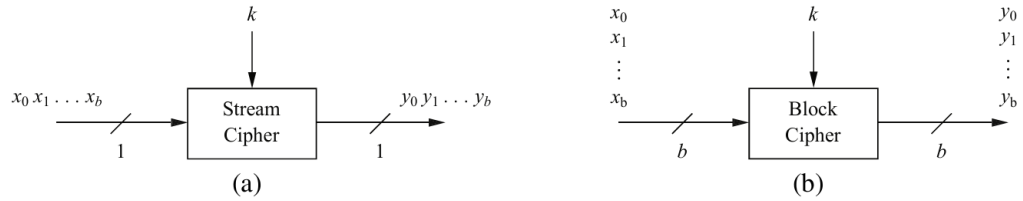


Figure 1.5: A comparison of *Stream* and *Block ciphers* [1]. The diagram illustrates the fundamental difference between the two symmetric-key encryption methods: a Stream Cipher encrypts data one bit at a time, transforming the input bit stream ($x_0, x_1, \dots$) into an output stream ($y_0, y_1, \dots$); while a Block Cipher encrypts data by processing it in fixed-size blocks of b bits, transforming an input block ($x_0, \dots, x_b$) into an output block ($y_0, \dots, y_b$).

For the sake of completeness, it is worth noting that within symmetric cryptography, there are two main categories of algorithms: *stream ciphers* and *block ciphers*. The main difference lies in how data is encrypted (or decrypted) using the shared secret key (see Fig.[1.5]).

Stream Ciphers

Stream ciphers encrypt message data one bit at a time. This requires first generating a "keystream" (s_i) from a secret key (k), which is then combined with the bits of

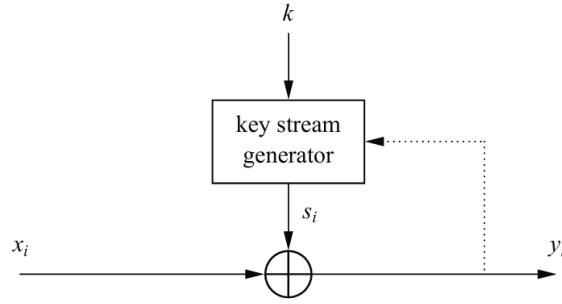


Figure 1.6: The basic schema for a synchronous and asynchronous stream cipher[1]. The key stream generator, initialized with a secret key (k), produces a keystream (s_i). This is used with a bitwise XOR operation ($\oplus$) to encrypt the plaintext bit (x_i) into the ciphertext bit (y_i). The dotted line represents the feedback loop in an asynchronous stream cipher, where the keystream depends on a previous ciphertext bit, while in a synchronous stream cipher, it depends only on the initial secret key.

the cleartext to be encrypted. This operation is typically an XOR (Section 1.4). These systems can be classified as *Synchronous* if the keystream depends only on the initial key (k), or *Asynchronous* if the keystream depends on both the key (k) and a previously produced piece of the ciphertext (y_{i-1}) (see Fig.[1.6]). The main characteristics of these schemes are their simplicity and speed, which makes them suitable for devices with limited resources.

Block Ciphers

Unlike stream ciphers, which work at the single-bit level, block ciphers, as shown in Fig.[1.7], encrypt the plaintext by processing entire blocks of bits at a time using the same key. The main advantage of this approach is that the bits of the plaintext are not encrypted independently of each other. The encryption of each bit within a block depends on all the other bits in that same block. This characteristic makes these cryptographic schemes extremely efficient, as it allows for the use of two primitive operations to build strong cryptographic algorithms[1]:

- *Confusion*: A cryptographic operation used to obscure the relationship between the key (k) and the ciphertext (y_i). It is generally achieved by substituting certain bits of the ciphertext with others.
- *Diffusion*: Similar to confusion, this is a cryptographic operation used to obscure the relationship between the key (k) and the ciphertext (y_i). However, in this case, a single symbol of the plaintext is used to influence many other



Figure 1.7: Basic scheme of a Block Cipher and the Data Encryption Standard (DES) algorithm [1]. The left side shows the general principle of a block cipher, which processes a fixed-size block of plaintext and a secret key to produce a block of ciphertext. The right side illustrates the specific parameters of the DES algorithm, which encrypts a 64-bit plaintext block using a 56-bit key to generate a 64-bit ciphertext block.

symbols of the ciphertext. This is generally achieved through a bit permutation.

In particular, these two operations minimize the probability that an attacker, upon intercepting the encrypted message, can reconstruct the original message by exploiting any periodicity of the bits of the key or the original message that might be reflected in the ciphertext. It is worth remembering, however, that even in this case, the fundamental encryption operation is usually the XOR.

Today, the most well-known modern block ciphers are the Advanced Encryption Standard (AES) and the Data Encryption Standard (DES) (see Fig. [1.7]), which use a standard block length of 128 bits (16 bytes) and 64 bits (8 bytes), respectively. While stream ciphers were previously favored for their efficiency in terms of computational and memory resources, contemporary block algorithms have achieved a level of optimization that allows them to leverage modern hardware and software architectures, making them the most widely adopted solution for securing internet communications.

1.2.1 Key Distribution

The most fundamental hypothesis of symmetric cryptography is that Alice and Bob must both possess the exact same secret key. To achieve this, before they can communicate a message over an insecure channel, the key must first be securely transmitted through a separate, protected channel. This is essential to ensure that

a malicious third party, Eve, cannot also acquire the key and thus decrypt the secret message (y) once it's intercepted.

Various methods can be used to accomplish this. For instance, the key—which can be a simple numeric string—could be physically carried on a hard disk from Alice to Bob. However, a key point to remember for these systems is that once the key has been exchanged, as long as it remains confidential, it can be used for multiple subsequent communications between the parties involved.

Finally, unlike what one might expect, both the encryption and decryption algorithms are publicly known to encourage analysis by as many users as possible, who can then identify and correct any weaknesses. With these observations, the secure communication problem is reduced to two main challenges:

1. Transmitting the key secretly.
2. Storing the key securely.

This concept is summarized by *Kerckhoffs's Principle*, which all robust cryptographic systems should respect.

Definition 1.2.1 (Kerckhoffs's Principle [1]). A cryptographic system must be secure even if the attacker (Eve) knows all the details of the system (such as the encryption and decryption algorithms), with the exception of the secret key (k).

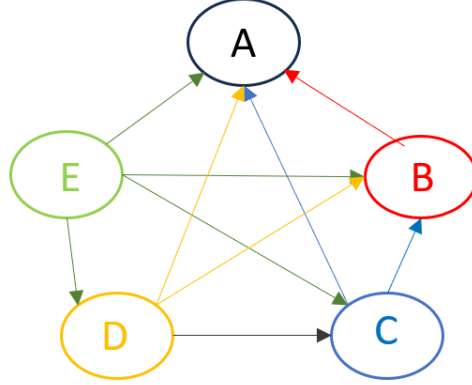


Figure 1.8: An illustration of the *Key Distribution Problem* in a Symmetric Cryptosystem. For a complete network graph with n nodes, the number of required secret key pairs is given by the formula: $C_2 = \frac{n(n-1)}{2}$. [2]

This figure shows an example with five nodes, requiring a total of 10 unique key pairs. This rapid increase in the number of keys to be managed and stored securely represents a significant challenge for symmetric cryptography, making it impractical for large-scale networks.

The common characteristics of all modern symmetric cryptography schemes are security and speed. This is particularly clear in algorithms like DES and 3DES. However, an intrinsic disadvantage of these schemes, related to the very nature of symmetric cryptography, arises from the basic assumption that the secret key must be shared before the actual communication between Alice and Bob using a different secure channel from the insecure one intended for the message. In fact, even if we manage to distribute the key securely (we will see that Quantum Key Distribution, or QKD, plays a decisive role here), as the number of users in a network increases, we have to manage a very large number of secret keys.

For instance, in a network with n users where each pair requires a unique secret key, it can be demonstrated that the total number of key pairs to manage is equal to $n(n-1)/2$. Consequently, each individual user must securely store $(n-1)$ keys (as shown in Fig.[1.8]). This scenario poses a significant problem even in a medium-sized network. For example, a company with approximately 1500 employees would need to manage over 1 million key pairs, which not only have to be securely generated but also distributed through equally secure channels to the network users.

1.2.2 The Random Nature of the Key

One of the most crucial aspects for making our symmetric cryptographic system secure is to ensure that the key shared between the parties involved (or the *keystream* derived from it) remains perfectly secret.

To be truly secure, a key must be completely random. In fact, if it has any pattern or regularity, an attacker can use different statistical techniques to decrypt the message, even if the key is not intercepted. Consequently, complete randomness ensures there are no predictable patterns for an attacker to exploit to break the cipher and read the secret message. A common solution to this, particularly in the context of *stream ciphers*, is to start with a small, shared secret initial key. This key is then used by a generator to produce a much longer and entirely random sequence of bits.

Today, there are three main types of such generators [1]:

1. **True Random Number Generators (TRNGs):** TRNGs produce completely random sequences of numbers based on various physical phenomena, such as the electronic noise in semiconductors or radioactive decay. It is important to remember that Quantum Random Number Generators (QRNGs) also fall into this category. As we will explore in Chapter 2, by leveraging the intrinsic principles of quantum mechanics, they can be considered a purer and more secure type of TRNG, and for this reason, they represent the future of modern cryptography.

2. **Pseudorandom Number Generators (PRNGs):** PRNGs generate sequences of numbers that are not purely random but are instead deterministic, meaning they can be calculated from a starting value called a *seed*. In general, many types of PRNGs have good statistical properties that allow us to create numerical sequences similar to those we would obtain from a TRNG, and for this reason, they are used today in many computer applications. However, there are possible attacks that can be performed by a potential attacker, which makes these generators insufficient in the cryptographic domain to guarantee that a system can be defined as having *Unconditional Security* (Def.1.2.3).
3. **Cryptographically Secure Pseudorandom Number Generators:** This is a particular class of PRNGs to which a fundamental characteristic is added: unpredictability. This essentially means that, assuming a given sequence of n consecutive bits of the key is known, no algorithm exists that can predict the next bit ($n + 1$) in polynomial time with a success probability greater than 50%. These types of generators are designed specifically for cryptographic applications and are created and optimized specifically for software or hardware implementation.

1.2.3 Unconditional Security and the One-Time Pad

We can now finally understand when such a cryptographic system can be defined as unconditionally secure.

Definition 1.2.2. *Unconditional Security* [1]. A cryptographic system is said to be unconditionally or theoretically secure when it cannot be broken even with unlimited computational resources.

The most important aspect of this definition is that it places no limits on an attacker's computing power. Consequently, to define such a system, one must necessarily assume that the attacker has access to the best modern technologies, such as a classical supercomputer or a quantum computer. This assumption is fundamental and makes the requirement for an unconditionally secure system a very difficult undertaking from both a theoretical and practical point of view. In fact, even if a cipher is practically unbreakable today, it cannot be defined as unconditionally secure, because an attacker with infinite resources could simply test all possible keys in a single step.

Despite this, a very simple, unconditionally secure cipher has been theorized. It is called the *One-Time Pad* (OTP), and it requires the three conditions listed below.

Definition 1.2.3. *One-Time Pad* (OTP)[1]. A stream cipher is a One-Time Pad if:

1. The key stream $(s_0, s_1, s_2, \dots)$ is generated by TRNG.
2. Only the parties involved in the communication (Alice and Bob) know this key stream (Kerckhoffs's Principle).
3. Each bit of the key stream (s_i) is used only once.

Under these conditions, the OTP is rigorously defined as an unconditionally secure cipher. This was formally demonstrated by Shannon in his paper "Communication Theory of Secrecy Systems" [3].

Limitations and Disadvantages of the OTP

It is important to note that, despite its unconditional security, the OTP is generally not used for common applications today. Its main limitations can be directly identified in its definition and can be summarized in three fundamental requirements:

- **Requirement 1 (TRNG):** The OTP requires the use of a true random number generator, which is generally not present in standard commercial PCs.
- **Requirement 2 (Secure Transport):** Alice and Bob must exchange the key stream through an intrinsically secure channel. Obviously, in reality, a 100% secure system cannot exist, as even if the key stream is encoded on a physical medium, we would still need to deliver it through a person who is theoretically trustworthy. A system of this type is impractical for most applications.
- **Requirement 3 (Single Use):** This can be identified as the main limitation for a practical application of the OTP. Requiring that each bit of the key can be used only once implicitly means that the key must have the same length as the message to be encrypted. For practical applications, this requirement is particularly limiting.

For example, if Alice and Bob want to exchange a 1 GB file, they will need a 1 GB key, and once it has been used, they will have to discard it and exchange another one of the same size. This is obviously not feasible for rapid and secure communication today.

In conclusion, although the OTP is rarely used in practice for these reasons, it remains a very important theoretical system. It allows us to understand the fundamental requirement to strive for when designing cryptographic systems that, at their theoretical limit, can be defined as unconditionally secure.

The reason for this security essentially base on the property of true randomness.

Since each bit of the key is chosen from a truly random source, and it is used to encrypt a bit of cleartext—using the XOR operation—only once, the resulting ciphertext is statistically independent from the plaintext. This means that for any given ciphertext, every possible plaintext of the same length is equally probable.

Therefore, by successfully performing the elementary XOR operation (Section 1.4) between the bits of the cleartext (x_i) and the truly random bits of the key (s_i), the resulting ciphertext (y_i) cannot be broken by a potential attacker. Even exploiting infinite computational capacity, he (she) cannot distinguish the correct plaintext from all other possibilities, as they are all equally valid.

1.3 Principles of Asymmetric Cryptography

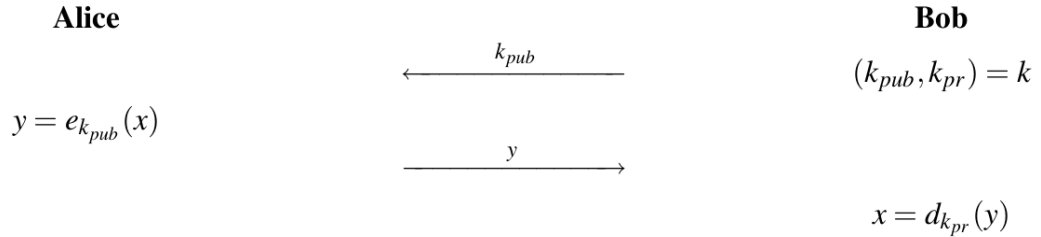


Figure 1.9: The basic schema for Public-Key (asymmetric) encryption[1]. The recipient, Bob, has a key pair composed of a public key (k_{pub}) and a private key (k_{pr}). Alice encrypts her plaintext (x) using Bob's public key, creating the ciphertext (y). Bob then decrypts the message using his unique private key. The public key can be openly shared, while the private key remains secret, eliminating the need for a secure key distribution channel.

In the previous subsection, we began to familiarize ourselves with the world of cryptography and, in particular, we introduced the basic principles of symmetric cryptography. It is worth remembering that, although this type of cryptography allows for the creation of particularly fast and efficient algorithms, for modern applications, the problem of distributing the secret key in a network of n users is a limiting factor as n increases. It is for this very reason that, starting from the second half of the 20th century, Diffie, Hellman, and Merkle proposed a new and revolutionary cryptographic scheme.

In this new scheme, it is no longer necessary for the key held by the person encrypting the message (Alice) to remain secret. Instead, it can be revealed to anyone, thus becoming a *public key*. Conversely, only the key of the person for

whom the message is intended (Bob), which is now used exclusively for decryption, remains a *private key*. A cryptographic system of this type is therefore called *Asymmetric*, since, unlike the symmetric one, the key used by Alice and the one used by Bob are no longer the same (Fig.[1.9]).

The key idea behind this scheme is that Bob holds a key pair (k) composed of two closely related parts: the public key (k_{pub}), known to everyone, and a private key (k_{pr}), exclusive to the decryption procedure. The principle underlying the operation of asymmetric cryptography is based on a particular type of function called a *one-way function*.

Definition 1.3.1 (One-Way Function) A function $f(x)$ is a one-way function [1]. if:

1. $y = f(x)$ is computationally easy to calculate, meaning it must be executable in at least polynomial time.
2. $x = f^{-1}(y)$ is computationally infeasible to calculate, meaning that this operation is so intensive that it would require an extremely long time to perform, even with the best available algorithms.

Today, two mathematical problems commonly rely on the difficulty of one-way functions in asymmetric cryptography[1]:

- *The integer factorization problem:* Given a very large integer, it is extremely difficult to find its prime factors. Conversely, given two very large prime numbers, it is easy to calculate their product.
- *The discrete logarithm problem:* Mathematically, given the finite cyclic group $\mathbb{Z}_p^*$ of order $p-1$, a primitive element α in $\mathbb{Z}_p^*$ and another element $\beta \in \mathbb{Z}_p^*$, the problem consists of determining the integer $1 \leq x \leq p-1$ such that: $\alpha^x \equiv \beta \pmod{p}$.

It follows that, while there is a vast number of symmetric algorithms, the asymmetric algorithms of practical relevance can be classified into three main families, each secured by specific mathematical problem:

- **Integer factorization schemes:** These are based on the integer factorization problem introduced above. The most important example among them is the RSA algorithm.
- **Discrete logarithm schemes:** These are based on the discrete logarithm problem. The best-known example is the Diffie-Hellman Key Exchange.

- **Elliptic curve (EC) schemes:** These are derived from a generalization of the discrete logarithm problem. The most popular examples are the Elliptic Curve Diffie-Hellman (ECDH) and the Elliptic Curve Digital Signature Algorithm (ECDSA).

Practical Aspects of Public-Key Cryptography

As we have discussed, the primary advantage of public-key cryptography over symmetric-key systems is its simplified key management. However, a significant drawback is that public-key algorithms are much slower. They not only require longer keys, which inherently slow down the encryption process, but the entire data encryption operation is also exceptionally time-consuming.

To give a sense of the scale, this process can be hundreds to thousands of times slower than the same operation performed by symmetric algorithms. Consequently, public-key cryptography is not typically used for encrypting the actual data, but its most important application is found in modern hybrid protocols, such as SSL/TLS and IPsec, which are designed to create secure web connections. More precisely, these schemes combine the high speed of symmetric ciphers for data encryption with the benefits of public-key algorithms for secure key exchange and non-repudiation. For example, these protocols could first use a public-key algorithm (like RSA) to securely exchange a symmetric key (for example, an AES key) between Alice and Bob, and then after this initial communication, both parties use the secret key to encrypt and decrypt subsequent messages with a much faster cipher.

Finally, it is crucial to address another major challenge in asymmetric cryptography: public key authenticity. How can we be sure that the public key we receive from Bob truly belongs to him, and not to a malicious attacker like Eve? This is the basis of a "man-in-the-middle" attack, where Eve intercepts Alice's request for Bob's public key and sends her own key instead. Consequently, because Alice encrypts her message with the wrong key, Eve can intercept and decrypt the message with her private key, compromising the communication.

The solution to this problem is the use of digital certificates, which can be seen as a digital credential that cryptographically binds a specific public key to the identity of its owner.

Chapter 2

Quantum Cryptography

2.1 Quantum 2.0

From the First to the Second Quantum Revolution

For over three-quarters of a century, quantum mechanics has been at the center of heated discussions among physicists and philosophers debating the very interpretation of its probabilistic nature. Innumerable scientific experiments and theoretical demonstrations carried out since the 20th century have aimed to understand and accept a new way of conceiving the measurement process. This new perspective is no longer tied to a mere observation of physical phenomena but rather to the intrinsic uncertainty that characterizes the nature of human knowledge itself.

The profound impact of these debates, which primarily took place in the first half of the 1900s, led to what we now call the First Quantum Revolution (or Quantum 1.0). Its effects have never truly faded, as many purely quantum phenomena continue to challenge researchers and force us to constantly question our interpretation of reality.

However, in recent decades, the scientific community's approach has undergone a radical transformation. We have now accepted that certain quantum phenomena behave in "bizarre" ways—that is, in a manner inconsistent with classical physics—and have instead sought to harness this "strangeness" on a purely technological level. This has led to the advent of a new era, the Second Quantum Revolution (or Quantum 2.0)[4], which has materialized through the rapid development of revolutionary technologies based on the most powerful features of quantum mechanics, such as the no-cloning theorem, quantum entanglement, and teleportation.

To date, the results of this second great quantum revolution have manifested in two main areas: quantum computing and quantum cryptography. The advent

of the Quantum Computer has radically changed the direction of scientific and technological progress. In particular, through a complete redefinition of the concept of the bit (now a qubit) and, consequently, of classical computation, the quantum computer stands out for its superior computational capabilities compared to classical computers. The consequences of such a powerful machine will impact various fields of our lives, but in this context, we will focus on its implications for modern cryptography.

By exploring the principles underlying quantum cryptography and studying real-world experimental setups for the implementation of quantum cryptographic protocols, it is possible to fully understand the revolutionary impact that the Second Quantum Revolution will have on our digital lives.

2.1.1 The Impact of Quantum Computing on Modern Cryptography

Today, the most common strategy for securing our digital information relies on a specific type of function known as a one-way function (see Def.[1.3.1]). These functions are computationally easy to execute in one direction but are extremely difficult to reverse.

A particularly widespread asymmetric cryptographic system that uses a one-way function based on the integer factorization problem is the RSA algorithm[5]. When two large prime numbers are chosen, it is computationally simple to calculate their product. However, it is computationally infeasible to perform the inverse operation, i.e., to find the prime numbers that, when multiplied together, produce it.

The Shor's Algorithm It's important to understand that the RSA algorithm's effectiveness is based on using prime numbers so large that factoring their product would take a prohibitively long time on a classical computer. For instance, in modern RSA implementations, this time can exceed a quadrillion years, which is far longer than the age of the universe [5].

The key point here is that the security of these cryptographic systems—the very foundation of all modern technology—isn't based on the impossibility of decryption. They do not constitute an unconditionally secure system (see Def.[1.2.3]). Instead, their security depends on the computational limitations of today's classical computers. It would take so long to break systems protected with RSA that they are considered practically unbreakable.

Mathematically, this is evident because the best-known classical algorithm for factoring a semiprime integer has a runtime of $\mathcal{O}(\exp(d^{1/3}))$, where d is the number of decimal digits of the integer [5].

Currently, the best factorization result achieved with a classical computer is for a number with a length of about 250 decimal digits and required approximately 2700 CPU-core years [6]. Therefore, these cryptographic systems still appear more than efficient for our digital security. However, with the advent of quantum computers, it has been demonstrated that a particular "hybrid" algorithm (one that uses both classical and quantum principles) called Shor's algorithm can solve the integer factorization problem on which RSA is based. Instead of exponential time, it does so in polynomial time of $\mathcal{O}(d^3)$.

This implies that with the advent of quantum computers the entire RSA system on which much of modern cryptography is build (including many secure financial transactions on the Internet) would be compromised.

Another crucial aspect to consider is that the fragility of today's cryptographic systems isn't just a current problem; it's a future threat. If a malicious party were to intercept and store encrypted messages today, even without the immediate technological or economic resources to break them, they could simply wait for the development of a sufficiently powerful quantum computer. It follows that, in the near future the security and confidentiality of today's secret messages will be in danger.

2.1.2 The Advent of Quantum and Post-Quantum Cryptography

This scenario became explicit to the global community in 2016, when the U.S. National Security Agency (NSA) issued a public warning to the world's cryptographic community [7]. In this document, the NSA expressed concern about the potential threat that quantum computers pose to cybersecurity at both governmental and commercial levels. From that moment on, government agencies and the scientific community have responded in two main ways to find adequate countermeasures:

- To develop and standardize new *Post-Quantum Cryptography* (PQC) algorithms.
- To leverage the potential of *Quantum Key Distribution* (QKD) to create unconditionally secure cryptographic systems.

In the approach known as Post-Quantum Cryptography, the idea is to develop new classical cryptographic systems that are robust against attacks from quantum algorithms. A possible limitation of this approach is that keep using purely classical algorithms to counteract the ever-increasing computational power of current quantum computers does not completely solve the problem. There could be quantum (or even classical) algorithms that are still unknown and could violate the security

of these new systems in the future. In other words, the solution offered by post-quantum cryptography it does not guarantee unconditional security (Def.1.2.3)

In contrast, an approach that leverages the potential of Quantum Key Distribution (QKD) offers a complete and definitive solution for re-establishing global cybersecurity. This new approach is no longer based solely on "purely classical" cryptographic systems, but on new hybrid systems that utilize a coexistence of the principles of classical and quantum mechanics. However, the general situation remains quite complex today. While specific types of QKD protocols have already been demonstrated in real-world applications, their practical implementation is still challenging and their performance is often limited.

Consequently, the approach generally favored today involves implementing more practical QKD protocols. To achieve higher performance, these protocols often rely on a certain degree of trust in the technologies used, which are not perfectly secure. This compromises the "unconditional" security promised by QKD, making the systems vulnerable to various kinds of attacks [8]. Furthermore, even if we were able to overcome this trade-off between security and performance, we would still face a critical limitation: the relationship between the key generation rate and distance. This factor strongly limits the applicability of current point-to-point QKD protocols to the typical cryptographic applications we use every day, a topic that we will explore in detail in Chapter 5.

2.2 Quantum Cryptography

2.2.1 Origin of Quantum Cryptography

While the first ideas for using quantum mechanics in cryptology date back to the early 1970s, when Wiesner proposed leveraging its principles To create secure banknotes that cannot be copied, and the first scientific paper on quantum cryptography was published in 1982, the birth of quantum cryptography is generally dated to 1984. This was the year Charles Bennett and Gilles Brassard published a revolutionary scientific article introducing one of the first and most well-known quantum cryptographic schemes: the BB84 QKD protocol [9].

The paper stems from the consideration that, if asymmetric cryptographic systems (like RSA) are or will be vulnerable to attacks from quantum computers, the only solution is to create unconditionally secure systems (see Def.[1.2.3]) using symmetric (or private-key) cryptographic schemes.

The revolutionary idea of the two scientists was not to transmit the secret message directly with a quantum signal, but to use such a signal to distribute the secret cryptographic key to the two parties. In this way, the scheme proposed by Bennett and Brassard became the first example of Quantum Key Distribution (QKD), and as we will see, it stands apart from asymmetric cryptographic systems because it can resist attacks from a quantum computer.

For clarity and completeness, it's worth noting that a secret key distributed via quantum cryptography could, in principle, be used in one of two ways:

1. As a one-time key in cryptographic schemes like the OTP (*one-time pad*).
2. As a shared secret key in a standard symmetric cryptographic protocol like DES or AES.

Let us discuss these two cases separately [2].

In an OTP scheme, the idea is to use the key as a one-time mask (Section (1.2.3)). However, the main problem with this scheme is that, although it is truly unbreakable under the condition of a completely random key, it is practically inconvenient to have to transmit new secret key bits at the same rate as the message bits are sent due to the low achievable key rates.

In this second application scenario, Alice and Bob do not require a One-Time Pad, but rather a secret key to be used with a private-key cryptosystem such as the Advanced Encryption Standard (AES). The key's distribution could be managed using asymmetric cryptosystems like Diffie-Hellman or RSA; however, these are susceptible to decryption by a future quantum computer. Quantum Key Distribution (QKD) offers a solution by providing a secure method for key exchange.

As it has been demonstrated that quantum distribution protocols are immune to quantum computer attacks, QKD provides a robust alternative [8]. Furthermore, the key rate provided by QKD is sufficient for use with AES and similar algorithms where the key is significantly shorter than the message, making this approach a more practical and viable solution for contemporary cryptographic needs than the One-Time Pad, which demands a key of equal length to the message [10].

2.2.2 Introduction to DV-QKD

Today there are two main families of QKD protocols: *Discrete-Variable QKD* (DV-QKD) and *Continuous-Variable QKD* (CV-QKD). Although in this work we will focus exclusively on a particular type of DV-QKD protocol known as *Twin-Field Quantum Key Distribution* (TF-QKD), we will first briefly highlight the main theoretical and technological differences between them [11].

DV-QKD, based on discrete variable systems described by a finite-dimensional Hilbert space, was the first to be introduced with the well-known BB84 protocol [9] and relies on single-photon detectors and sources. In contrast, CV-QKD, introduced about 15 years later, is based on coherent detectors and sources and uses an infinite-dimensional Hilbert space, making it currently less sensitive to the losses and noise of today's technologies.

Preliminary Notions

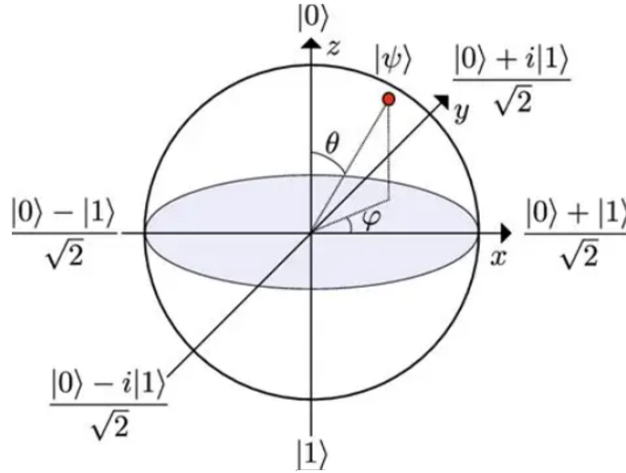


Figure 2.1: Representation of the Bloch Sphere [12]

Discrete-variable protocols can be considered the first and most fundamental form

of Quantum Key Distribution. They therefore represent the most suitable starting point for familiarizing oneself with quantum cryptography and for introducing some preliminary concepts, such as the qubit and its representation, which will be useful to us later.

The Qubit as the Fundamental Unit of Quantum Information. The qubit, or quantum bit, represents the fundamental unit of quantum information. It is defined as any two-level system that can exist in a superposition of two quantum states, typically denoted as $|0\rangle$ and $|1\rangle$. Unlike a classical bit, which can only be either 0 or 1, a qubit can exist in a special condition where it is both 0 and 1 at the same time. This property is central to both quantum computation and quantum cryptography and stems from a core principle of quantum mechanics called the *Principle of Quantum Superposition* [13].

Mathematically, a qubit is represented as a two-dimensional vector within a *finite-dimensional Hilbert space* [14]. The basis vectors of this space are:

$$|0\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \quad \text{and} \quad |1\rangle = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

Consequently, any state a qubit can assume can always be written as a linear superposition of these two basis states:

$$|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$$

where α and β are complex numbers such that $|\alpha|^2 + |\beta|^2 = 1$.

Geometrically, the qubit is usually represented as a vector on the Bloch Sphere (Fig.[2.1]). This geometric representation helps us understand the concept of a quantum superposition state by using the spherical coordinates θ and ϕ to define the qubit's state [5].

- When $\theta = 0$ or $\theta = \pi$, the vector is at the poles. These positions represent the two non-superposition pure basis states: $|0\rangle$ (at the North Pole, $\theta = 0$) or $|1\rangle$ (at the South Pole, $\theta = \pi$).
- When $\theta = \pi/2$, the vector is on the equator of the sphere, representing a pure superposition state. For example, the states $|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ and $|-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)$ lie on the equator, along the X axis. Similarly, the states $|+i\rangle = \frac{1}{\sqrt{2}}(|0\rangle + i|1\rangle)$ and $|-i\rangle = \frac{1}{\sqrt{2}}(|0\rangle - i|1\rangle)$ are also on the equator, along the Y axis.

- For any other value of θ ($0 < \theta < \pi$), the vector is in an intermediate position on the sphere's surface, representing a specific pure superposition state. An example is the state $\frac{1}{\sqrt{2}}(|0\rangle + e^{i\pi/4}|1\rangle)$.

Mutually Unbiased Bases (MUB). The basis vectors $|0\rangle$ and $|1\rangle$ that were introduced are also called eigenstates of the Pauli matrix σ_z , and this matrix is what is commonly called the Z basis in QKD. Similarly, the states $|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ and $|-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)$ are called eigenstates of the Pauli matrix σ_x , which is called the X basis. Finally, the states $|+i\rangle = \frac{1}{\sqrt{2}}(|0\rangle + i|1\rangle)$ and $|-i\rangle = \frac{1}{\sqrt{2}}(|0\rangle - i|1\rangle)$ are eigenstates of the Pauli matrix σ_y and form what is called the Y basis. Where the Pauli matrices can be written as:

$$\sigma_z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad \sigma_x = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}, \quad \sigma_y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix}$$

Definition 2.2.1. Formally, two orthogonal bases of a d -dimensional Hilbert space, for example $\{|\psi_1\rangle, \dots, |\psi_d\rangle\}$ and $\{|\phi_1\rangle, \dots, |\phi_d\rangle\}$, are mutually unbiased (MUB) if $|\langle\psi_i|\phi_j\rangle|^2 = 1/d$ for all i and j . This means that the key property of MUBs is that if a state from one basis (e.g., a state from the Z basis) is measured in another basis (e.g., the X basis), we will get any of the states from the latter basis with an equal probability ($1/d$). Therefore, the result we obtain is not deterministic, but probabilistic.

NOTE This property is precisely the foundation of the security of all DV-QKD protocols. In fact, the most used bases in different DV-QKD protocols (like the well-known BB84) are the Z and X bases. Since they are mutually unbiased, if an eavesdropper (Eve) measures a qubit using the wrong basis (for example, by measuring a qubit encoded in the Z basis with the X basis), her measurement will randomly alter the state of the qubit. Consequently, when Bob performs his measurement with the same basis as Alice (for example, the Z basis), he will obtain a value that will not correspond to the one Alice sent. Finally, the errors introduced by this interception can be detected by Alice and Bob when they compare a small sample of their key, allowing them to confirm the presence of an eavesdropper.

Physical Interpretation and Encoding with Photons. As we have previously introduced, the qubit is defined as any two-level system that can exist in a superposition of two quantum states. This means there is no single way to implement them, but they can be prepared in various physical systems: from the

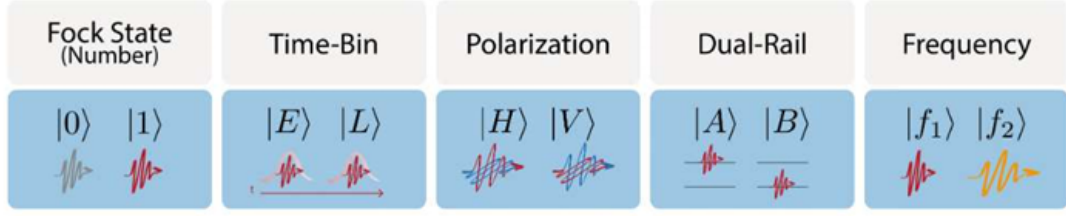


Figure 2.2: Different degrees of freedom for encoding quantum information using single photons [15]. The figure illustrates how a qubit can be represented by various physical properties of a photon, including its Fock state (number of photons), Time-Bin (arrival time), Polarization (horizontal or vertical), Dual-Rail (path), and Frequency. These properties are manipulated to encode the logical states $|0\rangle$ and $|1\rangle$ for quantum communication.

optical degrees of freedom of single photons, to the spin states of electrons, to the electrodynamic states of a superconducting circuit.

Despite this, the current technological approach is to use stationary systems (like a single atom) to implement what are called *quantum nodes* and then exchanged quantum information between these distant nodes through optical communication channels (free-space or optical fibers) in which photons are transmitted to encode and transport the information.

In summary, at a technological level, two types of qubits are essentially identified [15]:

- **Stationary qubits** at the separated nodes, which act as quantum memories to store the quantum information.
- **Flying qubits**, which are realized with photons to distribute the quantum information between the various nodes of the network.

NOTE. It is important to emphasize that for QKD, the only requirement is to allow the parties involved in the communication (Alice and Bob) to share the quantum information in which they have encoded the bit of the future shared secret key. Consequently, in these types of schemes, there is no need to store quantum information in devices such as the quantum memories (stationary qubits) of the network nodes. This is why in all the QKD protocols we will analyze, only photons are used for the realization of flying qubits.

Having now understood that photons are the optimal choice in QKD schemes for sharing quantum information between Alice and Bob, let's explore how they can encode this quantum information by exploiting the various degrees of freedom of

these light particles (Fig.2.2). There are several degrees of freedom that can be used for this purpose. Among the most common in DV-QKD schemes are: polarization (*Polarization encoding*), phase (*Phase Encoding*), amplitude (*Fock-state encoding*), timing (*Time-bin encoding*), spatial modes (*Dual-rail encoding*), and frequency (*Frequency encoding*). For this work, we will specifically introduce *Polarization encoding* with the BB84 protocol first, and then *Phase Encoding* with the TF-QKD protocol.

Let's begin by delving into these two types of photon encoding.

1) In Polarization encoding, the idea is to associate the basis states of a qubit with specific polarizations that can characterize a photonic particle. This approach is very widespread today in numerous DV-QKD schemes because these quantities are easily identifiable and measurable. For example, regarding the computational basis states previously introduced, which we have named $|0\rangle$ and $|1\rangle$, it is possible to associate:

- A **Horizontal Polarization** (H) for the state $|0\rangle$.
- A **Vertical Polarization** (V) for the state $|1\rangle$.

However, as we have previously understood, a qubit can also exist in a generic superposition state of the type $|\psi\rangle = \alpha|0\rangle + \beta|1\rangle$. Consequently, using the previous association, it is now perfectly legitimate to write this mixed state as $|\psi\rangle = \alpha|H\rangle + \beta|V\rangle$. For example, the quantum superposition state $|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$ will now be represented by a 45° diagonal polarization. This is what we call the diagonal state $|D\rangle = \frac{1}{\sqrt{2}}(|H\rangle + |V\rangle)$. Analogously to what happened on the $\{|0\rangle, |1\rangle\}$ basis, it represents a balanced superposition state of the basis states ($|H\rangle$ and $|V\rangle$). This means that when we measure a photon with an $|H\rangle$ or $|V\rangle$ polarization, we will have a 50% probability of finding a horizontal polarization and a 50% probability of finding a vertical polarization. Similarly, the state $|-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)$ will now be represented by a -45° antidiagonal polarization, which we will call the antidiagonal state $|A\rangle = \frac{1}{\sqrt{2}}(|H\rangle - |V\rangle)$. Finally, it follows that in this case, the Hadamard basis (or the X basis) $\{|+\rangle$ and $|-\rangle\}$ will be written as $\{|D\rangle$ and $|A\rangle\}$ [13].

To recap:

Table 2.1: Polarization Encoding of a Qubit

	Basis states	Equal superposition	General state
Qubits	$ 0\rangle, 1\rangle$	$ +\rangle$	$\alpha 0\rangle + \beta 1\rangle$
Polarized photons	$ H\rangle, V\rangle$	$ D\rangle$	$\alpha H\rangle + \beta V\rangle$

From a purely practical point of view, QKD schemes (like BB84) that use this

type of encoding measure the polarization of the photons in which the quantum information is encoded (associated with the classical one) through two instruments:

- Polarizing filters: non-isotropical optical filters that, depending on their orientation, only allow photons with a polarization compatible with their axis to pass.
- A photon detector (like a photodiode) that counts how many photons have passed through the polarizing filter.

2) In Phase Encoding, the idea is to encode the basis states of a qubit into the relative phase of a single photon (see Table 2.2). The main advantage of this type of encoding is its robustness to loss and noise in the communication channel.

NOTE As we will see in more detail with the TF-QKD protocol (Chapter 4), this encoding is used in more complex interferometric schemes. In these schemes, Alice and Bob do not send photons directly to each other, but rather to a third, central station (Charlie) through two distinct optical paths. The quantum information is encoded in the relative phase that Charlie measures between the two photons received from Alice and Bob. This measurement is performed using an interferometric apparatus (e.g., a Mach-Zehnder interferometer) which converts the relative phase difference into a measurable intensity difference [15].

In this case, the Computational Basis (or Z Basis) can be rewritten using two states that represent a phase difference of 0 and π . So:

- State $|0\rangle$: relative phase = 0
- State $|1\rangle$: relative phase = π

While the Hadamard basis (or X basis) is composed of two superposition states that have a phase difference of $\pi/2$ with respect to the Z basis. So:

- State $|+\rangle$: relative phase = $\pi/2$
- State $|-\rangle$: relative phase = $3\pi/2$

To summarize, we have:

Table 2.2: Phase Encoding of a Qubit.

Basis	State	Relative Phase
Computational (Z) Basis	$ 0\rangle$	0
	$ 1\rangle$	π
Hadamard (X) Basis	$ +\rangle$	$\pi/2$
	$ -\rangle$	$3\pi/2$

2.2.3 Generic Structure of a Prepare-and-Measure QKD Protocol

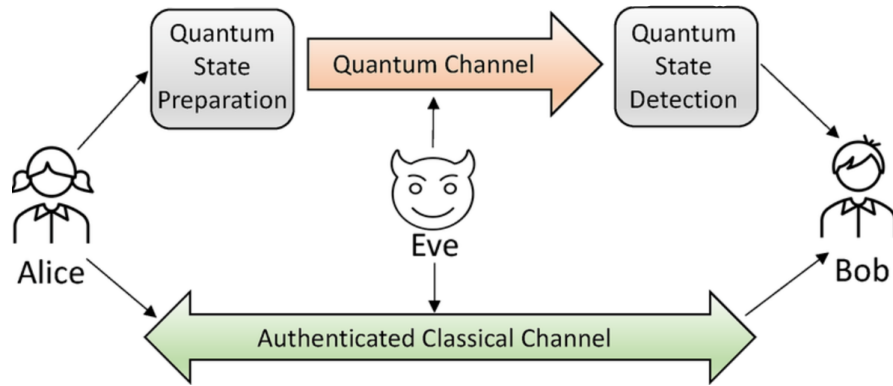


Figure 2.3: Basic schema of Quantum Key Distribution (QKD)[2]. Alice prepares quantum states and sends them to Bob via a quantum channel. Bob performs quantum state detection. An eavesdropper, Eve, may attempt to intercept the quantum channel. Crucially, an authenticated classical channel is also used between Alice and Bob for tasks like basis reconciliation and error correction, ensuring the security of the shared key even in the presence of Eve.

Today, there are multiple QKD protocols that differ not only in the physical quantities used to encode the bits of the key to be distributed, but also in the underlying principles of quantum mechanics and different mathematical security proofs. However, in any *Prepare-and-Measure* type QKD protocol ([Fig.2.3]), we can always identify two distinct and successive moments [8]:

1. Quantum communication
2. Classical post-processing

Quantum Communication

The initial phase of the protocol, referred to as quantum communication, follows a sequence of steps:

1. **Encoding:** The sender, Alice, transforms a random classical bit, α , into a quantum state. This state is selected from a set of non-orthogonal states, a critical feature for the protocol's security. This is achieved by manipulating a physical property of a photon to encode the information.
2. **Transmission:** These quantum states are then sent across a public quantum channel, such as an optical fiber, which is susceptible to interception by a third party, Eve.
3. **Measurement:** The receiver, Bob, measures the incoming quantum signals using a procedure complementary to Alice's encoding process, independently yielding his own random classical bit, β .

This procedure is repeated until Alice and Bob have collected a sufficient amount of raw data.

Ideal Scenario In an ideal communication, all the technologies used (source, channel, detectors, etc.) are free of errors and there is a completely secure quantum channel with no eavesdropping from Eve. When both Alice and Bob employ the same basis for encoding and measurement, their respective classical variables, α and β , will exhibit a perfect correlation or anti-correlation [8]. Consequently, this exchanged data can then be used directly to extract a final bit string that is completely identical and secret from an eavesdropper. This secret key can then be used to encrypt their future communication, even over an insecure classical channel.

Real-World Scenario In contrast to the ideal case, real-world implementations of Quantum Key Distribution (QKD) protocols are inherently limited by imperfections. Communication channels (free-space or optical fiber) are susceptible to noise, such as bit-flip or phase-flip errors, which inevitably introduce errors during the quantum communication phase [12]. Furthermore, the performance of the detectors can impact the final secret key rate (SKR). For example, due to limited *detection efficiency*, some photons may not be detected, or due to *dark counts*, false clicks can occur in the absence of an incoming photon [16].

Moreover, in a real-world scenario, we must also consider the possibility of an

eavesdropper, like Eve, performing an attack to gain information about the exchanged data. While various attacks exist, the simplest is when Eve acts as a *man-in-the-middle*, trying to intercept the quantum signal sent by Alice to Bob and measure the quantum state to discover the classical information (bit) α that Alice has encoded [8].

The security of this phase is fundamentally guaranteed by the No-Cloning Theorem, a core principle of quantum mechanics. This theorem states that an arbitrary quantum state cannot be perfectly duplicated [5]. As a result, Eve cannot simply make a copy of the quantum states sent by Alice without changing them. Instead, she can only attempt to measure them to gain information. However, because Alice and Bob use non-orthogonal measurement bases (Section 2.2.1), Eve does not know which basis to use. This uncertainty forces her to introduce errors into the quantum signals, which Alice and Bob can detect, thereby revealing her presence.

Consequently, the string of bits that Alice and Bob exchange in a real-world setting is not completely identical or secret due to these various errors. The exchanged data cannot be used directly as a private key; instead, it must be considered "raw data" to be processed in an **additional classical post-processing phase** to extract the final private shared key [16]. Furthermore, they must also find a way to determine if an eavesdropper has attempted to gain information about this key. To do this, they must essentially determine how much information Eve has obtained during her attack and consequently determine whether the communication is secure.

QBER in Quantum Key Distribution. To estimate the impact of all these non-ideal factors and to detect the presence of Eve, the concept of the Quantum Bit Error Rate (QBER) is introduced.

The Quantum Bit Error Rate (QBER) is a fundamental parameter used in any QKD protocol as a performance indicator. Essentially, it quantifies the ratio of erroneous bits to the total number of bits exchanged between Alice and Bob and is defined as [13]:

$$QBER = \frac{\text{Number of Wrong Bits}}{\text{Total Number of Received Bits}}$$

The probability of an error occurring is primarily related to the characteristics of the specific technology chosen for the QKD implementation in a real-world setting. This includes the communication channel's properties, such as noise and attenuation, as well as the characteristics of the components used, such as detector efficiency and dark counts.

In addition to these factors, any attack from an eavesdropper, like Eve, will introduce extra errors, causing the observed QBER to rise significantly. This

increase serves as a critical security alarm for Alice and Bob, signaling a potential compromise of the communication. If the estimated QBER after the quantum communication is higher than the value that characterizes the specific setup, it indicates that an eavesdropper is potentially active in the channel. In such a case, the communication must be declared insecure because the private key is potentially known by an unauthorized third party.

Ultimately, the final QBER value is essential for determining the necessary amount of post-processing required to distill a perfectly secure and identical key. In fact, various algorithms can be used for this purpose and the specific parameters chosen for these protocols depend on the estimated QBER for the communication under consideration.

Based on those considerations, we can conclude that the QBER is a comprehensive metric whose value is a complex function of intrinsic system imperfections, detector noise, environmental factors, and the specific rules of the QKD protocol being implemented. Therefore, understanding these various contributions is essential for designing robust and secure QKD systems [12]:

- **Intrinsic Bit-Flip Errors ($QBER_{intrinsic}$):** These are errors where a bit's value is flipped due to external noise or potential attacks. This type of error is protocol-dependent and reflects imperfections in the preparation or measurement steps.
- **Detector-Related Errors ($QBER_{detector}$):** These errors depend on the physical properties of the detectors, such as dark counts (detections that occur in the absence of an incoming photon) and background photons (detections caused by environmental light sources). These effects are particularly relevant in long-distance transmissions where the signal becomes weak.
- **Protocol-Specific Contributions:** The overall QBER is also shaped by the specific QKD protocol in use, as each has unique characteristics that affect how errors are counted and managed.

Classical Post-Processing

Now, let's analyze the main steps Alice and Bob must take in a real-world scenario. Because the data they share during quantum communication is in a "raw" form, only a portion of it will ultimately be used to create the final key.

The post-processing procedure can be divided into four distinct steps:

1. **Sifting:** The first routine Alice and Bob perform is called sifting. This is the

phase in which they publicly announce the sequences of bases they independently chose during quantum communication. They subsequently discard the bits shared in all those time windows where their chosen bases did not match.

The sifting phase is based on a particular application of *Heisenberg's Uncertainty Principle* according to which they can reconstruct the same bit with a probability of one only when they both choose the same basis. Therefore, they can share a common secret key only by using this data.

This same principle also justifies why it is possible to detect Eve's presence during the communication phase using non-mutually orthogonal states (Def.2.2.1): it can be used to demonstrate that the probability of a single measurement by Eve producing a detectable error for Alice and Bob is equal to 25%.

2. **Parameter Estimation and Error Characterization:** Following the sifting phase, Alice and Bob perform an intermediate step in which they publicly reveal a portion of the shared raw data. This action is crucial for characterizing the channel's properties, such as noise, and for evaluating the security of the exchange.

Although, both a noisy channel and an interception attempt by Eve introduce errors into the shared bit string. By publicly comparing a small, pre-agreed subset of their bits, Alice and Bob can precisely estimate the rate of these errors (QBER). In particular, a low QBER indicates a secure channel where the raw data can be used for key generation, whereas a high error rate is a strong indicator of Eve's interference. More precisely, if the error rate overcome a predetermined threshold, they must immediately abort the QKD protocol, as the security of the key cannot be guaranteed.

3. **Error Correction:** In this phase, Alice and Bob run a classical post-processing algorithm to identify and eliminate any discrepancies in their shared bit string.

At the end of this step, they are highly confident that their two shared keys are identical. However, the keys are not yet completely secret from Eve. This is because the public comparison of a subset of bits during the previous parameter estimation phase has potentially provided to Eve a small amount of information about the raw key. This partial knowledge, although minimal, is enough to compromise the key's perfect secrecy.

4. **Privacy Amplification:** This is the final stage of the post-processing procedure, where Alice and Bob aim to generate a completely shared secret key. The core idea is to use a universal hash function to compress the unsecured

shared string obtained from the error correction step, which results in a shorter but entirely secret final key. The primary goal is to eliminate any bits of the raw shared key that Eve might possess. The more bits that are discarded, the greater the secrecy of the final key, but the shorter its length will be [14].

At the end of these four phases, Alice and Bob obtain, through the QKD protocol, ideally **an identical and secret shared key** that can then be used to protect future messages they send through the insecure channel.

2.2.4 Overview of DV-QKD protocols

In this chapter, the idea is to provide a general overview of some of the most important DV-QKD protocols today, trying to understand their basic operating principles and the differences between them [8].

Prepare-and-Measure DV-QKD

One of the most well-known families of DV-QKD protocols called: *Prepare-and-Measure*, where the scheme is essentially the one illustrated in all our examples introduced so far, in which Alice (the sender) prepares optical quantum signals by encoding a random discrete classical variable (a bit) and then sends them to Bob (the receiver), who measures them to recover the information shared by Alice.

BB84. Historically, we have already mentioned that the first DV-QKD protocol to be introduced was the BB84 protocol [9] in which the discrete variable used for encoding the quantum information is the polarization of photons. In this protocol, to be precise, the two non-orthogonal bases used by Alice and Bob are the Z computational basis, encoded with a linear photon polarization and written as $\{|H\rangle, |V\rangle\}$ (where H is horizontal polarization and V is vertical polarization), and the Hadamard X basis, which is instead associated with its diagonal polarization, $(\{|D\rangle, |A\rangle\})$, where D is diagonal polarization and A is antidiagonal polarization.

The security of this protocol is based on the assumption that Alice and Bob use this set of non-orthogonal bases (Section 2.2.1) to exploit the *Heisenberg Uncertainty Principle* and the *No-Cloning Theorem* [5].

B92. Another particularly important protocol was proposed by Bennett a few years later in 1992 and is known as the B92 [17]. It can be considered the simplest

QKD protocol to date as it demonstrates that QKD can also be realized **with only two non-orthogonal quantum states** (e.g., $|\psi_0\rangle = |0\rangle, |\psi_1\rangle = |+\rangle$).

The security of this protocol, similarly to what happens in BB84, is based on the fact that since these states are not orthogonal, they return a non-zero scalar product that introduces an intrinsic uncertainty in the measurements and allows an attacker like Eve to be detected with a non-zero probability.

Generally speaking, however, it is shown that the performance of the basic scheme of the original B92 protocol is not as good as that of BB84. This is particularly linked to the fact that the use of only two linearly independent states makes it possible for an eavesdropper to perform a particular type of quantum measurement called an *Unambiguous State Discrimination* (USD [18]) measurement on the quantum states prepared by Alice, which makes the B92 very dependent on the losses and noise of the channel used.

Entanglement-based DV-QKD

Another important family of QKD protocols, called *Entanglement-based DV-QKD*. This family exploits one of the most fundamental principles of quantum mechanics: entanglement [15].

E91. The first scheme that for the first time exploited the purely quantum phenomenon of entanglement for cryptographic purposes was proposed in 1991 by Ekert with the E91 protocol [19]. Unlike the previous schemes introduced, in this case, the security of the protocol is based on a particular test, called the *Bell Test*, used to exclude the presence of an attacker (Eve).

For the implementation of this protocol, a more complex scheme is used in which there is a single source that emits pairs of entangled particles (e.g., polarized photons), described by a Bell state (e.g., the superposition state $|\Psi\rangle = \frac{1}{\sqrt{2}}(|01\rangle - |10\rangle)$). These particles are then separately sent to Alice and Bob, who each receive one of the two halves of each pair. In this case as well, Alice and Bob measure the received particles by randomly choosing from a set of three possible bases (chosen according to the Clauser, Horne, Shimony, and Holt (CHSH) test [20]) and perform all the steps described in Section (2.2.3) that characterize a generic QKD protocol.

The substantial difference compared to BB84 is, however, that in E91 all the temporal slots of the quantum communication in which Alice and Bob have chosen different bases are not simply discarded, but are used to verify the presence of Eve through the *Violation of Bell's Inequality* [21]. In particular, in the absence of an

attack, if Alice and Bob detect this violation, they can be sure that no attacker has compromised the entanglement and that, since their communication is secure, they can use the data exchanged in the temporal slots where they chose the same basis to distill a shared private key.

BBM92. The BBM92 protocol [22] can, in a sense, be conceived as a more sophisticated version of the BB84 protocol in which the principle of entanglement is exploited, just as is done in E91. However, it is shown that the BBM92 works more efficiently than the latter because it requires Alice and Bob to use only two different mutually unbiased measurement bases (Section 2.2.1), instead of the three bases of E91, which can be chosen in a similar way to those of the BB84 protocol.

In this case as well, Alice and Bob perform all the steps described in Section (2.2.3) that characterize a generic QKD protocol, and unlike the E91 protocol, the basic idea for its security is based on the fact that, when Eve becomes entangled with Alice and Bob's qubits, she inevitably causes errors in their measurements that can lead to her being detected. Consequently, in BBM92 there is no need for the legitimate parties to engage in a Bell test.

NOTE A QKD protocol with a Bell test provides a higher level of security because it allows for verifying that the measurement results are purely "quantum" and are not the result of a pre-programmed interaction or a hardware defect in Alice and Bob's devices that Eve could exploit to steal the secret key without being detected.

Device-Independent DV-QKD

All DV-QKD protocols that we have mentioned until now, base their security proofs on precise theorems and mathematical demonstrations, which in turn are based on specific assumptions.

For example, in traditional QKD protocols (like BB84 or BBM92 without a Bell test), it is assumed that Alice and Bob's devices are perfect, i.e.:

- Devices are assumed to generate the quantum states with perfect accuracy.
- They are expected to perform measurements in the correct bases.
- They must not allow a potential eavesdropper to exploit other degrees of freedom (such as spin or polarization) to gain information without being detected.

Today, however, even if rigorous security proofs exist, it is difficult to assume the existence of devices that satisfy these assumptions, and consequently, we must always assume that there is a non-zero probability that any real characteristic of

the devices used in the QKD protocol and not included in the security proof could compromise the security of the latter. Attacks that exploit these characteristics are known as *Side-Channel Attacks* [8].

DI-QKD. Device-Independent Protocols (DI-QKD) [23] were developed to solve a major problem in quantum cryptography. Unlike traditional methods that assume perfect devices, these protocols' security proofs make no assumptions about how the devices work internally. Relying on this, DI-QKD protocols are resistant to a wide range of attacks, including those that leverage implementation vulnerabilities or side-channels, which is why they are seen as the ultimate solution for quantum key security [24].

In particular, DI-QKD protocols ensure security by checking the input-output behavior of the devices during the protocol. This verifies that the devices are working correctly, eliminating the need for regular hardware inspections. The core of this process is *Bell's theorem* [21]: when two separate devices (that can't communicate with each other) receive random inputs, and their outputs show a distribution that violates Bell's inequality, it proves their outputs were not predetermined. This allows Alice and Bob to use these devices in a DI-QKD protocol to generate a secure, shared private key.

Despite their strong security, DI-QKD protocols are generally harder to implement than traditional QKD protocols like BB84. This is mostly because they rely on entangled photon pairs, which are often lost during long-distance transmission. These losses, caused by the attenuation of the optical fiber, reduce the rate of key generation in any QKD protocol. However, they become a critical issue for DI-QKD due to the *Detection Loophole* [24] a security vulnerability that arises from the inefficiency of the photon detectors themselves (see also Appendix ??).

More precisely, in a device-dependent protocol like BB84, detector inefficiencies just slow down the key generation, but Alice and Bob can still "post-select" the successful events where photons were detected to build their shared key. In a device-independent protocol, instead, if the detector efficiency (η) drops below a certain threshold [25], the non-detection events can be used by an eavesdropper to create a false Bell violation and, consequently, allow them to manipulate the devices without being detected. To be more precise, this loophole becomes a threat when the detector efficiency falls below the specific Eberhard limit of $2/3$ [26].

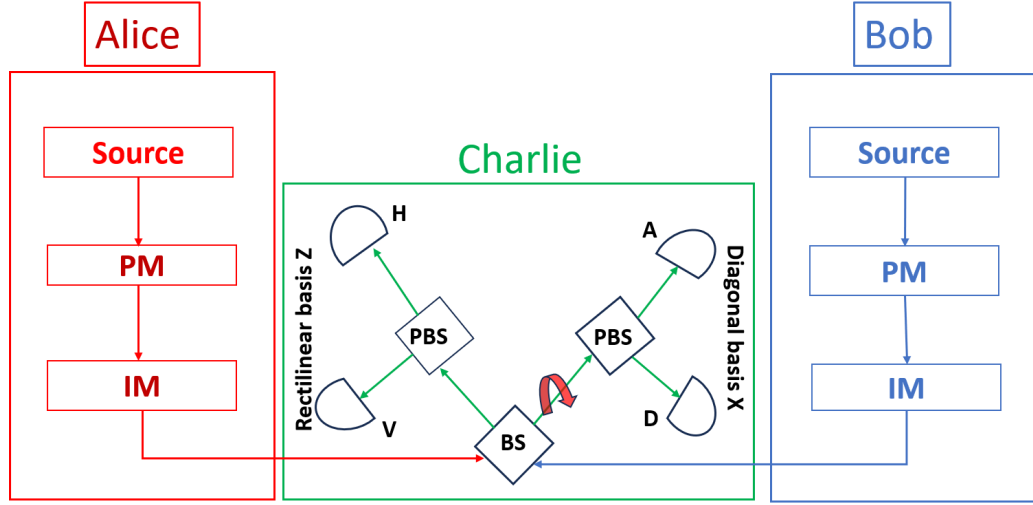


Figure 2.4: Basic schema of the Measurement-Device-Independent Quantum Key Distribution (MDI-QKD) protocol. Alice and Bob each prepare quantum states and send them to a central, untrusted node, Charlie. Charlie performs a Bell State Measurement (BSM) on the incoming signals and publicly announces the result. This protocol's security is guaranteed because the BSM result reveals a correlation between Alice's and Bob's bits without Charlie learning the actual values, making the protocol robust against *Detector Side-Channel attacks*.

MDI-QKD. A special and very important case of DI-QKD is the MDI-QKD (Measurement-Device-Independent QKD) protocol [27]. Its specific purpose is to eliminate vulnerabilities related to attacks on measurement devices (the detectors) [25]. However, it remains vulnerable to other types of attacks, such as source attacks (e.g., photon number splitting attacks), which require the use of additional techniques like "decoy-state" methods [28].

In this protocol, Alice and Bob do not exchange quantum signals directly. Instead, each sends their signal, over insecure channels, to an untrusted intermediate node who is potentially controlled by the attacker Eve, commonly called Charlie where the protocol's measurement devices are placed. Therefore, Charlie can then perform a simultaneous measurement on the quantum signals sent by the respective terminal nodes, known as a Bell State Measurement (BSM)[21].

This process "projects" the overall quantum state of the sent signals into one of the four possible Bell states. Consequently, At the end, Charlie detects and records the measurement result and announces it publicly. Let us indicate each result of his measement a value which we denote by γ .

The key of the protocol lies in how Alice and Bob interpret Charlie's result in

the subsequent post-processing phase, which is entirely classical. Depending on the value of γ , they can deduce whether their original bits were correlated or anti-correlated, without Charlie knowing the exact values.

- If $\gamma=0$ or $\gamma=3$, Alice and Bob know that their bits were identical (both '0' or both '1'). This pair of bits can be used for the future key.
- If $\gamma=1$ or $\gamma=2$, they know that their bits were different (one '0' and the other '1'). In this case one of the two participants (e.g., Alice) can simply invert their bit to make it identical to Bob's, and thus this bit pair can also be used.

In this way, Alice and Bob can ideally obtain a shared secret bit string (in a real-world scenario, the situation is different and they must perform additional steps, as reported in Section 2.2.3).

Overall, the security of this protocol lies precisely in the fact that Charlie, while knowing whether Alice's and Bob's bits are correlated or anti-correlated, does not know their specific values (e.g., he does not know if they are '00' or '11'). Thanks to this mechanism, the final shared key remains completely secret, even if Eve had total control of Charlie's node.

TF-QKD The MDI-QKD protocol represents a significant step towards securing networks with untrusted intermediate nodes. However, it is fundamentally limited by its inability to surpass the theoretical maximum for point-to-point QKD, known as the *PLOB bound* [8].

Recently, this limitation has been overcome by a more efficient protocol introduced in 2018 [29] called Twin-Field (TF-QKD), which can be seen as an "advanced variant" of MDI-QKD. The key innovation of TF-QKD is its ability to achieve a secret key rate that depends on the channel transmittance (η) as the square root of the transmittance ($\sim \sqrt{\eta}$). This represents a significant improvement over the linear dependence ($\sim \eta$) of previous protocols and makes it a fundamental resource for future long-distance QKD communications. We will analyze TF-QKD in detail in the following chapters, along with one of its most notable variants, the sending-or-not-sending (SNS) QKD protocol.

2.2.5 Photon Number Splitting attack and Decoy states

The security of traditional QKD protocols, such as the well-known BB84, is continually at risk due to vulnerabilities in real-world implementations that can be exploited by a potential attacker to gain information about the final key. While completely Device-Independent (DI) QKD protocols have not yet been fully realized, the state of the art is represented by modern QKD protocols like

Measurement-Device-Independent (MDI)-QKD (including its advanced variant, TF-QKD). These are designed to close side-channel attacks by being independent of the measurement devices. However, even these advanced protocols remain susceptible to eavesdropping attacks at the source.

A primary example of this vulnerability is the Photon Number Splitting (PNS) attack, which leverages a common technological limitation to extract information without being detected. The following sections will detail the nature of these imperfections, explain how they can be exploited to perform a PNS attack, and present the countermeasures developed to mitigate them, thereby enabling secure key distribution over practical distances.

Weak Coherent Laser Source and the PNS Attack

Generating perfectly indistinguishable single photons with high efficiency is a technologically complex and expensive process. For this reason, most Discrete-Variable (DV)-QKD protocols today rely on highly attenuated weak coherent laser sources [12]. This solution is simple and low-cost, but it introduces a critical vulnerability: these sources do not emit a single photon with every pulse. Instead, the number of photons, n , in a given pulse follows a Poisson distribution:

$$P(n) = \frac{e^{-\mu} \mu^n}{n!} \quad (2.1)$$

Here, μ is the average number of photons per pulse.

This means there is always a non-zero probability of generating multi-photon pulses, even when the protocol is designed for single-photon states. This imperfection is precisely what a Photon Number Splitting (PNS) attack [8] exploits.

In a PNS attack, an eavesdropper, Eve, performs a *non-destructive measurement* [8] to determine the number of photons in a given pulse. If a pulse contains more than one photon, she can steal the excess photons and forward the remainder to Bob. This compromises security because Eve obtains a perfect copy of the encoded information without introducing any detectable errors in the photon stream that reaches Bob. Since Bob cannot detect her presence, Eve can simply wait for Alice to publicly reveal her basis during the sifting phase, then measure her stolen photons to obtain complete information about the shared secret key.

Consequently, the presence of multi-photon pulses significantly reduces the final secure key generation rate. The effective key rate in the presence of a PNS attack can be expressed as [8]:

$$R_m = [\mathcal{P}(1 - r_{\text{PA}}) - H_2(e)]Q \quad (2.2)$$

Here, Q is the ratio of total bits in the raw key to total detected signals, $\mathcal{P}$ is the fraction of those detected signals that were single-photon pulses, r_{PA} is the Privacy Amplification rate, and $H_2(e)$ is the rate associated with error correction.

NOTE This formula highlights that the secure key rate depends directly on the fraction of single-photon pulses, which is unknown in a real-world implementation. [8].

The Decoy-State Method

As we've established, the use of weak coherent laser sources introduces a critical vulnerability: the presence of multi-photon pulses that can be exploited by a PNS attack. This imperfection forces Alice and Bob to discard a significant number of bits, dramatically reducing the final secure key rate. To overcome this, a countermeasure known as the *decoy-state method* was developed [14].

This approach is based on the security analysis from Improved Lo-Ma-Gottesman-Lo-Preskill (ILM-GLLP) [28], which demonstrated that a secure key rate could be obtained even with imperfect sources, provided that two crucial quantities were accurately determined: the gain (Y_n) and the Quantum Bit Error Rate (QBER) (e_n) for each pulse of n photons exchanged in the quantum communication. Among all possible values, the most critical are those for single-photon pulses (Y_1 and e_1), as only these pulses are considered truly secure.

In a practical QKD system, Alice and Bob can directly measure the total gain (Q_m) and total QBER (E_m) of the channel. These measured values are weighted averages of all detected pulses and are given by:

$$Q_m = \sum_{i=0}^{\infty} Y_i \exp(-\mu) \frac{\mu^i}{i!} \quad (2.3)$$

Here, Y_i is the probability that Bob correctly detects a pulse containing i photons sent by Alice.

$$E_m = \frac{1}{Q_m} \sum_{i=0}^{\infty} Y_i e_i \exp(-\mu) \frac{\mu^i}{i!} \quad (2.4)$$

The challenge lies in the fact that Alice and Bob do not know the individual values of Y_i and e_i for each photon number, making it impossible to calculate the secure values Y_1 and e_1 . This is where the decoy-state method provides a solution. Instead of using a single average photon number μ for all pulses, Alice intentionally sends a subset of her quantum signals with different, carefully chosen intensities. These pulses are the *decoy states*.

By using at least two different μ_i values (a "signal state" μ_1 for key generation and

one or more "decoy states" μ_2 for parameter estimation), Alice and Bob can create a system of linear equations based on the formulas above. Solving this system allows them to accurately determine the unknown yields (Y_i) and error rates (e_i) for each photon number, including the critical Y_1 and e_1 values. This provides a way to estimate the secure key rate and effectively filter out the insecure bits.

It is important to remember that these decoy states are used exclusively for parameter estimation and are never used for key distillation. The final key is generated only from the pulses with a μ_1 value specifically chosen to have a high probability of containing a single photon. The choice to send a decoy state or a signal state is made randomly and independently by Alice, so Eve cannot know which pulses are for key generation and which are for surveillance. The number of photons chosen for the decoy states is a critical design parameter, as it directly impacts Alice and Bob's ability to obtain the most relevant values of Y_1 and e_1 for their specific protocol.

NOTE Theoretically, the summation in Eq. (2.3) and Eq. (2.4) extends to infinity, which would, in principle, require an infinite number of decoy states to determine Y_1 and e_1 with perfect precision. However, practical DV-QKD experiments have demonstrated secure implementations using a finite number of decoy states.

Chapter 3

TWIN-FIELD QKD PROTOCOL

3.1 Rate-Distance Limits of Current QKD Protocols

Based on our discussion of the foundational principles of QKD, let's now address a critical challenge: integrating this technology into long-distance optical fiber networks.

A major obstacle for long-distance QKD is that quantum channel loss directly reduces the secret key rate (SKR), which is the fundamental metric for measuring a protocol's performance. In particular, the transmittance, η , of a quantum channel decreases with its length, L , according to the relation $\eta = 10^{-\alpha L/10}$, where $\alpha = 0.2$ dB/km is a typical fiber attenuation coefficient. This dependency means that for most single-photon QKD protocols, like BB84, the rate of photon detection is directly proportional to η . Consequently, the SKR for these traditional protocols decays exponentially with distance, even under ideal, noise-free conditions.

To fully understand and quantify the performance of QKD protocols, we need to consider both the ultimate theoretical limits and the practical constraints.

Secret Key Capacity (SKC)

First, let's consider the theoretical maximum. The Secret Key Capacity (SKC) is not a physical limitation, but a theoretical parameter that quantifies the ultimate maximum amount of secret information that can be transmitted securely over a given quantum channel [16]. It represents the highest possible secret key rate that Alice and Bob could ever hope to achieve, regardless of the protocol or technology

used. Although no current protocol can reach this limit, the SKC serves as a crucial reference point for the efficiency of different QKD schemes, setting the theoretical limit.

The Linear Limit

The linear limit, expressed as $R \leq -\log_2(1 - \eta)$ [30], is a fundamental constraint that applies specifically to point-to-point QKD protocols that do not use any intermediate measurement station. This formula, where R is the key rate and η is the channel transmittance, shows that the key rate of these protocols is inherently limited by the exponential decay of photons with distance.

The PLOB Bound

The Pirandola-Laurenza-Ottaviani-Banchi (PLOB) [8] bound is the most restrictive known rate-distance bound for all QKD protocols that do not use quantum repeaters. This bound represents the highest possible secret key rate achievable for repeaterless QKD. While protocols that use intermediate measurement stations, such as Measurement-Device-Independent (MDI)-QKD, can overcome the linear limit by mitigating some of the channel losses, they are still fundamentally constrained by the PLOB bound.

Unlike the linear limit, which only applies to point-to-point protocols, the PLOB bound provides a more general limit for all repeaterless schemes, whether they are point-to-point or employ intermediate measurement stations. Essentially, it represents the absolute performance ceiling for QKD without relying on technologically complex quantum repeaters.

3.2 TF-QKD: A Fundamental Resource for Next-Generation Long-Distance QKD

While today's fiber-based QKD systems have achieved a level of maturity that allows for their application in real-world telecommunication networks, the primary challenge remains the signal loss inherent in optical fibers.

A definitive solution could be the use of special devices called *Quantum Repeaters*, which could theoretically improve this scaling to $\eta^{1/(N+1)}$ with N intermediate nodes. Unfortunately, the technology required for these devices is still not practical. Consequently, due to this significant technological challenge, current implementations of the most important QKD protocols are limited to very short fiber distances (50–100 km) [31], well below the PLOB bound. This means that current QKD technology is characterized by a relatively low SKR, which is inadequate for typical

cryptographic activities, especially for inter-city communications that span several hundred kilometers.

A possible, but imperfect, solution for covering these inter-city distances would be to use QKD schemes based on trusted nodes. These are physically secure, intermediate stations that measure incoming quantum signals and then re-transmit new ones. The idea is to chain these nodes together to create a single long-distance link composed of a series of shorter, more manageable QKD links. However, this approach has a crucial limitation: the security of these schemes relies on the assumption that the intermediate node is fully secure [12].

NOTE While a trusted-node network can greatly extend the reach of quantum communications, these nodes are not true quantum repeaters because the incoming quantum information is "re-qualified" classically, effectively ceasing to be quantum.

We can therefore conclude that QKD schemes based on trusted nodes are not a definitive solution, as they do not exceed the SKC barrier. In 2018, a new QKD protocol called **Twin-Field QKD** (TF-QKD) was introduced. Although its operating scheme is similar to that of MDI-QKD, using an untrusted intermediate station (and thus also exceeding the linear limit), it has been shown to be capable of surpassing the PLOB bound [8].

3.2.1 TF-QKD Performance

The most important property of TF-QKD is that its secret key rate (SKR) exhibits the same distance dependence as a quantum repeater. This SKR scales as the square root of the channel transmittance ($\sqrt{\eta}$). We can demonstrate this property of the TF-QKD protocol using Fig.[3.1] [29], which provides a homogeneous comparison between theoretical limits (lines) and experimental results (symbols) for different quantum schemes implemented using optical fiber-based communication, all normalized to a standard optical fiber length L with an attenuation coefficient of $\alpha = 0.2 \text{ dB km}^{-1}$.

The theoretical limits shown are: I, MDI-QKD with decoy states [27]; II, QKD with decoy states [28]; III, single-photon QKD[8]; IV, SKC (indicated by a yellow dashed line) [31], ideal TF-QKD [29], and the limit one would have even with a single quantum repeater[12]. The experimental results reported are: discrete variable QKD [11], continuous variable QKD [32], and MDI-QKD, and they are indicated as squares, triangles, and circles, respectively, and are numbered in chronological order. Finally, a solid black line indicates the TF-QKD key rate obtained from a real experiment using a specific set of parameters, including: the dark count probability (P_{dc}); total detection efficiency (η_{det}); optical channel error rate (e_{opt}); and the error correction coefficient (f). In this way, we can compare this real rate

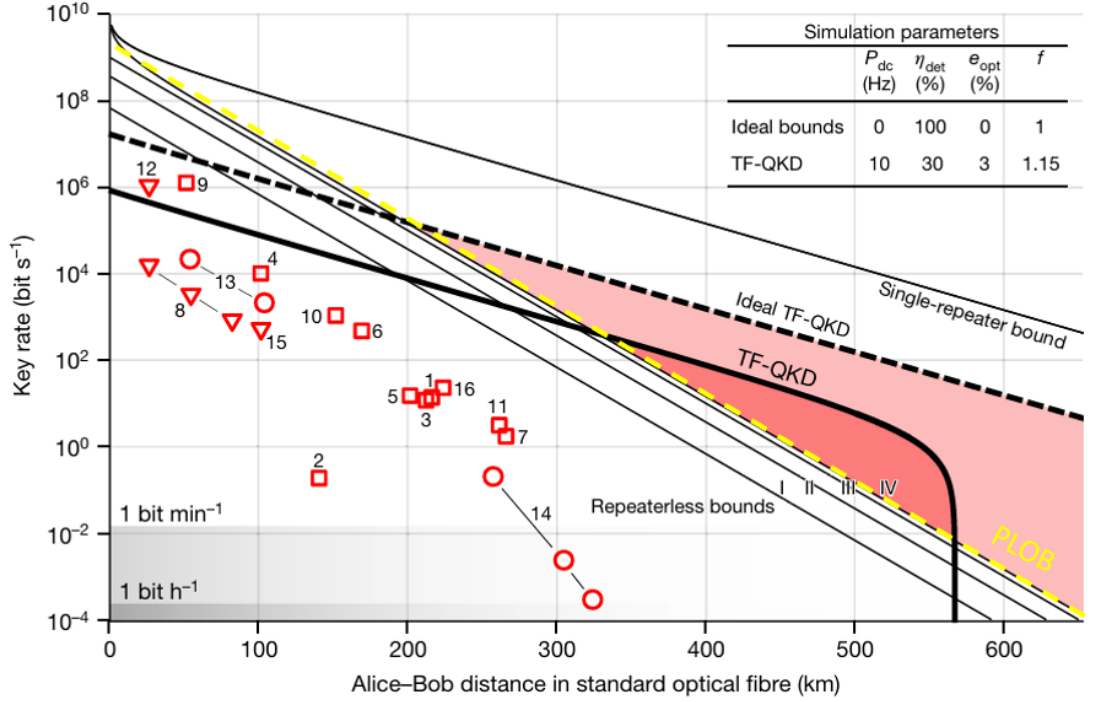


Figure 3.1: A homogeneous comparison between theoretical limits (lines) and experimental results (symbols) for different QKD schemes, reproduced from Lucamarini et al. [29]. The graph is normalized to a standard optical fiber with an attenuation coefficient of $\alpha = 0.2 \text{ dB km}^{-1}$. The theoretical limits shown are: I, MDI-QKD with decoy states [27]; II, QKD with decoy states [28]; III, single-photon QKD [8]; and IV, The maximum theoretical secret key capacity (SKC) is indicated by the PLOB bound (dashed yellow line) [8]; ideal TF-QKD [29], and the single quantum repeater limit [8]. The experimental results, indicated by squares, triangles, and circles, correspond to discrete variable QKD, CV-QKD, and MDI-QKD, respectively, numbered in corresponding chronological order. A solid black line represents the key rate from a real TF-QKD experiment, which, along with the ideal TF-QKD key rate (dashed black line), surpasses the repeater-less limits in the shaded area. This demonstrates that the TF-QKD key rate scales as the square root of channel transmittance (η), which is similar to a single quantum repeater, offering a significant advantage over other QKD protocols.

The inset table shows the key simulation parameters for the ideal bounds and the TF-QKD experiment: P_{dc} represents the detector dark count rate, η_{det} is the detector efficiency, e_{opt} is the optical misalignment error, and f is the error correction efficiency factor.

not only with the ideal rate (indicated by a dashed line), but above all with the SKC.

The most interesting results we can observe are:

1. From an experimental perspective, the key rates achieved by today's most efficient QKD protocols (indicated by red symbols) are characterized by rates that decrease as the communication length increases and never reach the maximum theoretical secret key capacity (SKC) indicated by the PLOB bound (dashed yellow line in Fig.3.1). This is the biggest limitation of the QKD schemes introduced so far and is strictly related to imperfections in the experimental apparatus used. These imperfections include, but are not limited to: the absence of a true single-photon source, noise in the communication channel, fiber attenuation, low efficiency and dark counts in detectors, and the presence of background photons.
2. Conversely, there are two areas (indicated in dark and light red) where both the real and ideal TF-QKD key rates exceed the PLOB bound. Specifically, ideal TF-QKD surpasses the repeater-less limits after 200 km, while real TF-QKD can surpass the ideal repeater-less limit after 340 km of optical fiber.
3. Furthermore, the way the TF-QKD key rate varies with distance is similar to that of a single quantum repeater, thus proving our previous statement: if the key rate for a generic conventional QKD protocol scales linearly with channel transmittance (η when $\eta \ll 1$), conversely, that of TF-QKD scales as $\sqrt{\eta}$.

We thus conclude that, although we have not provided a rigorous proof of the unconditional security of the TF-QKD protocol, it is clear from the preceding discussion that thanks to it, we are able to significantly improve the secret key rate compared to a traditional QKD protocol for the same quantum channel length or, equivalently, we can obtain the same SKR for a greater distance between the parties involved in the protocol.

NOTE While a formal proof of security is beyond the scope of this discussion, it's important to note that the robustness of the Twin-Field QKD protocol has been confirmed by several theoretical analyses. The protocol's architecture was designed to be resilient to a wide range of attacks as detailed in several works that have established its unconditional security guarantees. Among these, fundamental contributions are the studies by Curty et al.[33] and Yin et al.[31], which provided crucial validation.

Furthermore, in the next section, in which we will describe the operational scheme

of TF-QKD in more detail, we will demonstrate that the TF-QKD protocol is also a suitable solution for long-distance QKD from a technological point of view. This is another extremely important property because, unlike quantum repeater schemes, which promise to significantly exceed the SKR but are not yet practical, TF-QKD is based on a design that is implementable with current quantum technologies and that can guarantee manageable noise levels even when working on standard optical fibers hundreds of kilometers long.

3.3 Operating Scheme of TF-QKD

3.3.1 Basic Operating Principles

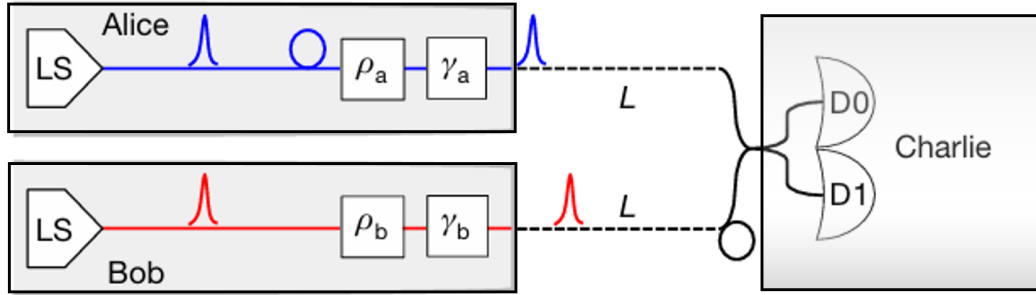


Figure 3.2: A basic operating scheme of the TF-QKD protocol [29]. Alice and Bob use distinct laser sources (LS) and modulators to prepare weak optical pulses with encoded phase information. They send these pulses through two independent optical fibers of length L to an untrusted central node, Charlie. At Charlie's station, the two pulses interfere on a beam splitter (BS), and the outcome is measured by two single-photon detectors (D0 and D1), which are then publicly announced.

The operational scheme of TF-QKD is conceptually very similar to the one we analyzed for the DV MDI-QKD protocol. It is important now to present the basic operating scheme of TF-QKD at a high level.

Essentially, the TF-QKD protocol works through the following steps, exploiting the basic operational schema shown in Fig.[3.2]:

1. Alice and Bob use two distinct attenuated light sources (laser source (LS) + variable optical attenuator (VOA)) to generate weak optical pulses.
2. They perform a precise phase and intensity modulation on these pulses to encode the classical information, i.e. the bits they want to share for the

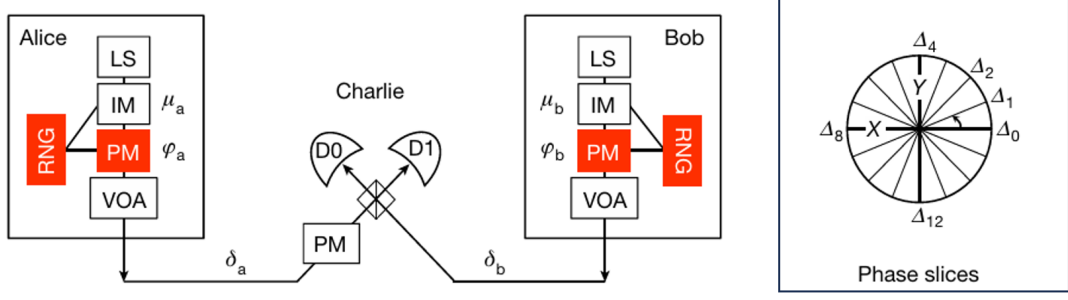


Figure 3.3: A detailed operating scheme of the Twin-Field Quantum Key Distribution (TF-QKD) protocol [29]. Alice and Bob each use a laser source (LS), intensity modulator (IM), and phase modulator (PM) to generate and encode weak coherent states with intensity (μ_a, μ_b) and phase information (φ_a, φ_b). These pulses are sent through separate optical fibers to an untrusted node, Charlie. At Charlie's station, the pulses interfere on a beam splitter (BS), and the outcome is measured by two detectors (D0 and D1). An additional phase modulator (PM) at Charlie's station is used for phase compensation, while the random number generators (RNG) at Alice and Bob's sites are used to choose phases and bases. The figure on the right illustrates the different *phase slices*, which correspond to distinct phase values on the Bloch sphere and they are used during the *Reconciliation Twin-Field Phase* step.

creation of the final key, based on the specific variant of the TF-QKD protocol being followed.

3. Both pulses are then sent through two distinct quantum channels to the untrusted intermediate station, Charlie.
4. In this central node, the interference of the pulses on a 50/50 beam splitter results in a detection event at only one of the two single-photon detectors (SPDs) if the phase of the initial pulses is coherent and remains coherent throughout propagation.
5. After the quantum communication is concluded, Charlie publicly announces a string that indicates which detector clicked for each time slot of the quantum communication.
6. Finally, from this public string, Alice and Bob can distill the final key through a classical post-processing phase.

3.3.2 Evolution toward the TF-QKD setup

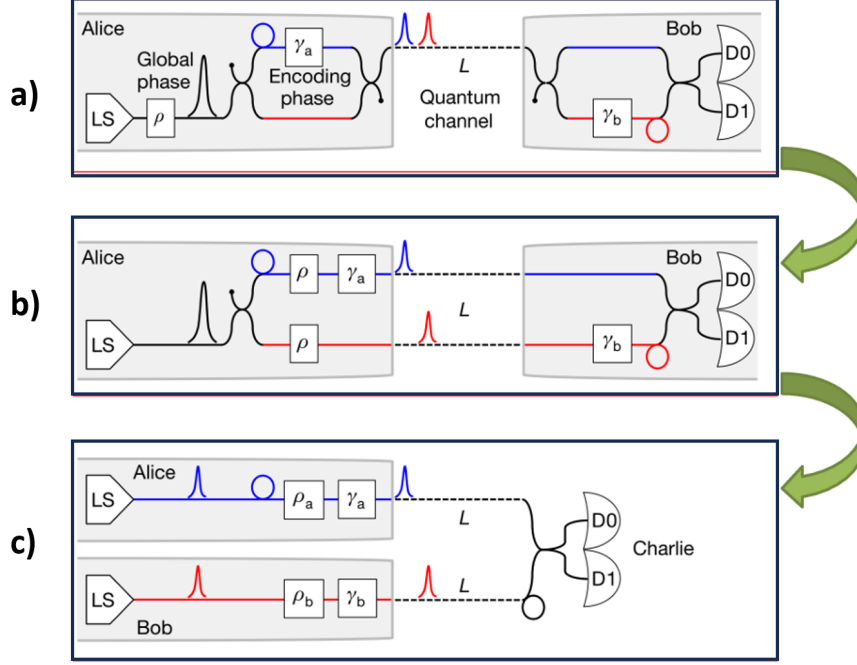


Figure 3.4: The evolution toward the final TF-QKD setup[29]. (a) A typical phase-based QKD protocol where a single source (LS) sends two pulses through an asymmetric Mach-Zehnder interferometer to Bob via a single quantum channel of length L . (b) An intermediate scheme where the two pulses travel on separate channels of the same length to Bob, conceptually leading to the TF-QKD idea of twin fields. (c) The final TF-QKD scheme where two separate sources from Alice and Bob send pulses over distinct channels to a central, untrusted node (Charlie). This setup doubles the effective distance between Alice and Bob to $L_{\text{tot}} = 2L$ for the same count rate, allowing for longer-distance quantum key exchange.

It is now interesting to understand how a typical phase-based QKD setup can naturally evolve into the TF-QKD protocol scheme just presented. This will allow us to derive the equation that describes the TF-QKD key rate directly from the standard formula used for the well-known BB84 protocol with decoy states.

Typical Phase-Based QKD Let's begin by describing a typical interferometric scheme used for the implementation of a generic phase-encoded QKD protocol (Fig.[3.4a]). In this case, a single light source (LS) emits optical pulses with a global random phase, ρ . Subsequently, the single pulse is split into two sub-pulses at the input of an asymmetric Mach-Zehnder interferometer, which creates a superposition

of states. A phase, γ_a , is encoded on one of these states. At this point, the two pulses are sent by Alice over a single quantum channel of length L to the receiving user, Bob, who also has an asymmetric Mach-Zehnder interferometer. Here, the incoming pulse is again split into a superposition of states, one of which is encoded with a phase, γ_b . The decoder distinguishes the pulses where a phase was encoded based on their arrival time. The interference of the pulses encoded with γ_a and γ_b enables the extraction of the raw key, which is then measured with detectors D_0 and D_1 .

Intermediate Scheme toward TF-QKD We can now slightly modify the previous scheme into a new setup (Fig.[3.4b]) where the two pulses now travel not on a single channel of length L , but on two separate channels of the same length.

In this case, both are encoded with the same phase ρ and undergo the respective phase shifts γ_a and γ_b at Alice's and Bob's locations, but the important thing is that they now represent what we have introduced in TF-QKD as *twin fields* that interfere at the central station (Charlie).

Beyond this difference, this new intermediate scheme is essentially equivalent, from the point of view of the type of coherent states generated and the classical information disclosed, to the previous scheme. Therefore, both the formula for the QKD key rate in the asymptotic scenario and the security principles can be derived directly from the latter.

The Final TF-QKD Scheme Finally, in Fig.[3.4c], we present the final scheme of the TF-QKD protocol. The main difference from previous schemes is the shift from a single sending node (Alice) to two separate sending nodes, Alice and Bob. Detection, on the other hand, is performed by a third central and *untrusted node*, Charlie, located at a distance L from both users. Consequently, the two users are separated by a total distance of $L_{tot} = 2L$.

It is crucial to understand the meaning of "untrusted node." This term implies that the node, in this case Charlie, could potentially be controlled by an attacker like Eve. Despite this, the protocol is designed to ensure that Charlie, while being able to know the information that passes through him (the interfering pulses), cannot in any way access or deduce the final key. The information he receives is essential for Alice and Bob to build the key, but it is not sufficient to compromise it. For this reason, the TF-QKD protocol, as an advanced variant of a Measurement-Device-Independent QKD (MDI-QKD) protocol (Section 2.4) [27], maintains its security even in this scenario [8].

In this scheme, Alice and Bob use two separate sources to emit optical pulses encoded with random global phases ρ_a and ρ_b . For these pulses to interfere

correctly on Charlie's beam splitter and allow for key generation, one fundamental condition must be met: the two emitted photons must be indistinguishable or "twin." This condition is achieved through a specific protocol routine known as *Coordinated Twin-Field Phase Randomization*, which differentiates TF-QKD from a traditional MDI-QKD protocol (see Section 3.4.2). The users distill the final key only after publicly revealing these global phases ($\rho = k + \phi$, where ϕ is the random phase and k is the encoding phase chosen by Alice and Bob) through the public announcement of the phase slices $\Delta k(a, b)$ [29].

It is important to emphasize that this public revelation of the two global phases represents a significant novelty compared to phase-based QKD protocols, where the global phase value is never revealed [16]. This constitutes a weak point in the security of the TF-QKD protocol, as it can leak information about the final key to a potential attacker like Eve. We will see later how some variants of the TF-QKD protocol, such as the SNS or CAL protocols, completely solve this security problem (Section 3.6).

In conclusion, the key feature of this final TF-QKD scheme is that, while in the previous typical phase-based QKD scheme the two pulses that encode the classical information co-propagate from Alice to Bob, in TF-QKD they travel from Alice and Bob toward the central node Charlie. This means that, for the same count rate, the pulses travel a halved distance to the detector, which allows the protocol to reach significantly greater transmission distances.

3.4 Properties of TF-QKD

3.4.1 TF-QKD vs MDI-QKD

Although the steps described must be tailored to the specific security proofs of the TF-QKD protocol variant, the general scheme shown in Fig.[3.4a] contains all the basic elements for implementing this protocol. According to this scheme, TF-QKD can be considered an evolution of the MDI-QKD protocol presented in Chapter 2.

This is true for at least three main reasons:

1. **The QKD Network Topology:** Both protocols utilize the same network topology, involving two trusted nodes (Alice and Bob) who want to share a private secret key, and an untrusted central third node (Charlie), who may be controlled by an attacker (Eve) and must therefore have no access to the exchanged key.
2. **The Protocol's Security:** As we saw for DV MDI-QKD, the basic security of the protocol is based on the fact that the untrusted node Charlie, being

able to only observe the result of this interference, can infer whether the secret bits chosen by Alice and Bob are the same (00 or 11) or different (01 or 10), but cannot learn their absolute values (0 or 1). Thus, in conclusion, even if Eve had control of this node, she would not be able to obtain any useful information about the final private key.

3. **Independence from the Detection Technology Used:** As in MDI-QKD, the security of TF-QKD does not depend on the specific measurement devices used.

However, there is an important difference between these two protocols that must be explained.

Interference Mechanism In the MDI-QKD protocol, the interference of photons sent by Alice and Bob at Charlie’s central station is followed by a Bell State Measurement (BSM) (Section 2.4) [21]. Consequently, to succeed, this measurement requires a coincidence count at Charlie’s detectors, confirming that both photons have arrived and have interfered as expected. The efficiency of this process is intrinsically limited by the probability that both photons survive during the communication and are detected together.

In contrast, in TF-QKD, the interference between the two photons sent by Alice and Bob produces a single-photon count in one of Charlie’s detectors. The protocol leverages the interference between the two twin fields to convert phase information into a single-photon detection event. This mechanism is much more robust and has a significantly higher probability of success.

It is precisely this fundamental change in the detection mechanism that allows TF-QKD to surpass the theoretical distance limit (known as the PLOB bound). Since the protocol’s success rate no longer depends on a two-photon coincidence event (whose probability is proportional to η^2 , where η is the channel transmittance), but on a single-photon event (whose probability is reduced only by the square root of the channel transmittance, i.e., $\sqrt{\eta}$), the TF-QKD protocol can operate over much greater distances than MDI-QKD.

3.4.2 Coordinated Twin-Field Phase Randomization

To ensure effective interference at Charlie’s central station, the total phase difference between Alice’s and Bob’s photons must fall within a well-defined range. Alice and Bob prepare their weak coherent states in the form $|e^{i(k_a+\phi_a)}\sqrt{\mu}\rangle_A$ and $|e^{i(k_b+\phi_b)}\sqrt{\mu}\rangle_B$ [31]. Here, k represents the phase (0 or π) chosen to encode the key bit, while ϕ_a and ϕ_b are the random global phases chosen from a continuous

interval $[0, 2\pi)$ [8].

The addition of these random global phases is a fundamental requirement for the protocol's security. Their exclusive role is to mask the key-bit encoding phases (k_a and k_b), preventing a potential eavesdropper like Eve from gaining knowledge about the encoded key bits.

The main problem that derives from this security solution is the fact that, due to the random nature of ϕ_a and ϕ_b , their total difference can fall into a range that doesn't allow for interference. This situation is defined as "No-Twin Fields." In such cases, the interference at Charlie is weak or non-existent, which prevents the correct measurement of the key bit and nullifies the entire process.

To solve this problem preventively, the TF-QKD protocol introduces an additional post-processing step known as *Coordinated Twin-Field Phase Randomization* [29]. Unlike traditional QKD where the global phase is kept secret, this procedure reveals the phases to a limited precision to guarantee effective interference. Instead of choosing a continuous phase, Alice and Bob divide the interval $[0, 2\pi)$ into M distinct sections. These sections each have a width of $\Delta\phi = 2\pi/M$ and are labeled $\Delta_k(a)$ and $\Delta_k(b)$ for Alice and Bob, respectively [31].

The key concept is that, to obtain effective interference, the random phases chosen by Alice and Bob must fall into the same predefined section. To identify these events, Alice and Bob publicly announce the strings corresponding to their chosen phase sections, $\Delta_k(a)$ and $\Delta_k(b)$, along with their preparation bases in the post-processing phase. This allows them to identify and discard all cases where $\Delta_k(a) \neq \Delta_k(b)$, and consequently, although their global phases are random, they must be coordinated to produce the necessary interference.

Intrinsic Quantum Bit Error Rate in TF-QKD It is interesting to note that, since the discriminating condition for identifying the twin field is that ϕ_a and ϕ_b fall into the same phase section ($\Delta_k(a) = \Delta_k(b) = 2k\pi/M$), then even in this case, where Alice and Bob do not discard the corresponding bit pair, these two random phases can still differ by at most a quantity less than $2\pi/M$ for a given pair of twin fields. It follows that the TF-QKD protocol will be characterized by an intrinsic QBER, E_M , due to the fact that the twin fields are close but not exactly identical.

It can be shown that, on average, this QBER will be [29]:

$$E_M = \frac{1}{\pi} \int_0^\pi \sin^2(Mt/2) dt = \frac{1}{2} - \frac{1}{2} \frac{\sin(M\pi)}{M\pi} \quad (3.1)$$

From Eq.3.1 we notice that E_M is minimised by choosing a higher number of intervals, thereby reducing the amplitude of each section. In an ideal case where

$M \rightarrow \infty$, this quantity tends to zero. On the other hand, the probability of matching two random phases in the same section scales with $1/M$.

From a practical point of view, it has been shown that an optimal value of M exists to maximize the performance of the TF-QKD protocol, achieving a value of $E_{M,opt} = 1.275\%$ with $M_{opt} = 16$ [29] in a real setup like the one previously described in Chapter 5. Nevertheless, as we will see, the *Matching Phase-Slices* operation is more complicated than this because in the real world, we must also consider that the phase is affected by noise and, consequently, assumes a non-constant value that drifts over time [16].

3.4.3 Derivation of TF-QKD Key Rate

Having shown a clear relationship between the typical scheme used for traditional phase-based QKD (Fig. 3.4a) and the main differences with respect to the traditional MDI-QKD scheme, all that remains is to derive the final formula for the TF-QKD key rate.

Let us remember that, in general, for the typical phase-based QKD experimental setup shown in Fig. 3.4a, the QKD key rate in the asymptotic scenario can be expressed according to the specific protocol being used. For example, if this setup uses the BB84 protocol with decoy states [28], the key rate can be written as [29]:

$$R_{\text{QKD}}(\mu, L) = \underline{Q}_{1|\mu,L} \left[1 - h(\bar{e}_{1|\mu,L}) \right] - f Q_{\mu,L} h(E_{\mu,L}) \quad (3.2)$$

where $R_{\text{QKD}}(\mu, L)$ is the QKD key rate in bits per pulse; $\underline{Q}_{1|\mu,L}$ and $\bar{e}_{1|\mu,L}$ are respectively the lower bound for the single-photon gain and the upper bound for the single-photon phase error rate, estimated using the decoy-state technique (see Section 3.5.3); $Q_{\mu,L}$ and $E_{\mu,L}$ are the measured total gain and QBER for a given laser pulse intensity μ sent on a channel of length L ; h is the binary entropy function; and f is a factor accounting for the efficiency of error correction.

NOTE: This equation explicitly shows how the key rate R is dependent on the total light intensity μ and, crucially, on the distance L between Alice and Bob.

Now, it is possible to express the TF-QKD key rate directly through Eq. (3.2) [29].

$$R_{\text{TF-QKD}}^{(\neg\rho)}(\mu, L) = \frac{d}{M} \left[R_{\text{QKD}} \left(\mu, \frac{L}{2} \right) \right]_{\oplus E_M} \quad (3.3)$$

The notation $\oplus E_M$ in the equation indicates the intrinsic QBER (E_M) of the TF-QKD protocol. The coefficient $1/M$ derives precisely from the selection of the phase slices $\Delta_k(a, b)$ described previously. Here, μ indicates the total intensity of

the optical pulses used, such that $\mu = \mu_a + \mu_b$, where μ_a and μ_b are the intensities of the pulses emitted by Alice and Bob. Finally, d indicates the duty cycle between the two possible modes that make up the initial communication between the users. As we will see later, they periodically alternate between classical and quantum communication phases to perform a phase realignment procedure, due to the non-ideal phase fluctuations that could destroy the Twin Field condition (Chapter 4) [8].

3.5 Secure TF-QKD Protocols: SNS and CAL

3.5.1 QKD Protocol Security Definitions

To fully understand the security of the TF-QKD protocol, let's begin by reviewing some essential definitions [8].

1. ϵ -Correct Protocol

A QKD protocol is defined as ϵ -correct if the probability that the final keys, S and S' , generated by Alice and Bob are not identical, is less than or equal to a value ϵ_{cor} . This is expressed as:

$$\Pr(S \neq S') \leq \epsilon_{\text{cor}} \quad (3.4)$$

This condition ensures that Alice and Bob will have identical keys with a very high probability, ensuring the reliability of the protocol's output.

2. ϵ -Secret Protocol

A QKD protocol is defined as ϵ -secret if, after Alice and Bob perform privacy amplification to produce two final secret strings of length l , the following condition holds for any attack by Eve on Alice's quantum state:

$$\frac{1}{2} \|\rho_{AE} - U_A \otimes \rho_E\|_1 \leq \epsilon_{\text{sec}} \quad (3.5)$$

In this inequality, ρ_{AE} represents the density operator of the combined system of Alice and Eve, U_A is the completely mixed state of Alice's string system of length l , and ρ_E is the density operator of Eve's system. This condition quantifies the secrecy of the final key. If the inequality is satisfied, it means that Eve's knowledge of the key is negligible because the state of the combined Alice-Eve system (ρ_{AE}) is nearly indistinguishable from a state where Alice's string is completely random and independent of Eve's information ($U_A \otimes \rho_E$).

3. ϵ -Secure Protocol

A protocol is defined as ϵ -secure if it simultaneously satisfies both the conditions for ϵ -correctness and ϵ -secrecy. The overall security parameter, ϵ , is bounded by the sum of the individual error probabilities:

$$\epsilon \leq \epsilon_{\text{cor}} + \epsilon_{\text{sec}}$$

This definition quantifies that the total error probability, which is a combination of the correctness error and the secrecy error, remains below an acceptable threshold.

3.5.2 Introduction to SNS and CAL Protocols

Today, numerous variants of the TF-QKD protocol have been introduced to overcome some critical issues of the original version. Let's analyze the most important ones for this work.

One of the main vulnerabilities of the original TF-QKD protocol, proposed by Lucamarini [29], is its weakness to a specific attack, known as the *beam-splitter collective attack*. This attack exploits a crucial step of the protocol where the terminal nodes, Alice and Bob, publicly declare their respective global phases to perform the phase reconciliation process (Section 3.4.2), a procedure that is fundamental for the assumptions underlying TF-QKD and for the subsequent construction of the shared private key.

More precisely, a key requirement of the original protocol is to keep only the bit pairs exchanged during communication slots where Alice and Bob have chosen the same basis and where their respective random global phases (ϕ_a and ϕ_b) fall within the same phase section ($\Delta_k(a) = \Delta_k(b) = 2k\pi/M$). The public disclosure of this phase section is the precise weakness that Eve exploits to perform a collective beam-splitter attack. The basic logic of this attack is as follows [8]:

1. Eve, who controls the untrusted central node (Charlie), knows that the total phase of each photon is a sum of two components: the secret key-bit encoding phase (k_a or k_b) and the random global phase (ϕ_a or ϕ_b).
2. Furthermore she knows that the interferometric measurement performed by Charlie does not depend on the individual encoding phase, but only on the total phase difference between the two transmitted photons, which can be expressed as:

$$\text{Total phase difference} = (k_a - k_b) + (\phi_a - \phi_b)$$

3. The vulnerability arises from the fact that, to identify successful sessions, Alice and Bob must publicly announce the phase section in which their random

phases fall. By intercepting this public information and identifying the cases in which both the random phases fall in the same section ($\Delta_k(a) = \Delta_k(b) = 2k\pi/M$), Eve discovers the difference in the global phases, $(\phi_a - \phi_b)$ of the quantum signals that are used by them to generate a bit of the secret key. Finally, since she can also measure the total phase difference of the photons arriving at Charlie's beam splitter, she can easily find the key's bit :

$$\text{Key bit} = (k_a - k_b) = (\text{Total phase difference}) - (\phi_a - \phi_b)$$

In this way, the public revelation of the global phase, while necessary for the protocol's operation, provides Eve with the critical information to discover the secret bit and compromise the final private key.

To solve this significant security problem, two variants of the TF-QKD protocol have recently been introduced, based on security proofs designed to be robust against coherent attacks. These variants are:

1. **SNS** (*Sending-or-Not-Sending*) [34]
2. **CAL** (*Curty-Azuma-Lo*) [35]

The most important characteristic of these protocols, which differentiates them from the original version and increases their security, is that the quantum communication between Alice and Bob is divided into two types of windows:

- **Signal Windows:** The user sends an appropriately modulated quantum state to exchange a potential bit of the future shared key.
- **Decoy Windows:** The user sends a *decoy* state used exclusively to accurately estimate certain parameters of the quantum communication through the channel and thus detect the possible presence of an attacker like Eve.

To this end, the SNS and CAL protocols employ two different and complementary groups of coherent states [31].

Remark: Coherent State

A coherent state is a quantum state that does not have an exactly defined number of photons (as in Fock states), but whose number of photons follows a Poissonian distribution around an average value $|\alpha|^2$. Specifically, such a state can be expressed in the form [8]:

$$|\alpha\rangle = e^{-\frac{1}{2}|\alpha|^2} \sum_{n=0}^{\infty} \frac{\alpha^n}{\sqrt{n!}} |n\rangle \quad (3.6)$$

NOTE A coherent state is an infinite superposition of Fock states, each corresponding to a different number of photons (0, 1, 2, ... n).

The two different groups of coherent states employed in these two TF-QKD protocol variants are [16]: *phase-mixtures states* in which the photon's phase is random, and *phase-defined states* in which the photon's phase is well-defined.

To fully understand the strategies adopted by the SNS and CAL protocols and how they are able to solve the original TF-QKD vulnerability to a beam-splitter attack, it is essential to delve into the properties of the two types of coherent states.

1. *Phase-mixtures*

Phase-mixtures are coherent states characterized by a random phase and an average intensity of μ . To an untrusted third party like Charlie or an eavesdropper like Eve, these states look like a statistical mixture of photon number states $|n\rangle$ that obey a Poissonian distribution. Consequently, the probability of detecting n photons in a state with intensity μ is defined by the following equation:

$$p_\mu(n) = e^{-\mu} \frac{\mu^n}{n!} \quad (3.7)$$

Given the weak intensity (μ), these mixtures are for the most part composed of either a vacuum state (zero photons) or a single-photon state. As a result, these states are associated with the *Z basis* ($|0\rangle, |1\rangle$). Consequently, their security proofs and the corresponding measurements are based on counting photons: detecting zero photons aligns with the $|0\rangle$ eigenstate, whereas detecting a single photon corresponds to the $|1\rangle$ eigenstate.

2. *Phase-defined states*

In contrast, Phase-defined states are coherent states whose components have a well-defined relative phase. This intrinsic property makes them suitable for an interferometric measurement, where the outcome is determined by comparing this internal phase relationship to an external reference. For this reason, this type of measurement requires a global reference phase to be established, either before or after communication with Charlie.

These states are defined as a superposition of the vacuum state ($|0\rangle$) and the single-photon state ($|1\rangle$). Consequently, we can say that they refer to the *X basis* ($|+\rangle, |-\rangle$), which is composed of two states distinguished precisely by their relative phase component.

- In the $|+\rangle$ state, the relative phase between the $|0\rangle$ and $|1\rangle$ components is zero.

$$|+\rangle = \frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$$

- In the $|-\rangle$ state the relative phase between the $|0\rangle$ and $|1\rangle$ components is π radians.

$$|-\rangle = \frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)$$

The respective security proofs are based on this precise coherent superposition of the vacuum state ($|0\rangle$) and the single-photon state ($|1\rangle$).

Comparison between SNS and CAL Protocols

Both the SNS [34] and CAL [35] protocols solve a potential security problem of the original TF-QKD by using both phase-defined states and randomized phase states. This approach creates two distinct windows in the communication (the Z-windows and the X-windows), which avoids the need to publish the global phase in every time window. In this way, the vulnerability to collective beam-splitter attacks is significantly reduced.

Although they are based on the same "security solution", these two protocols operate in a reciprocal manner. This reciprocity depends on which basis they exploit for the two main operations performed in TF-QKD protocol: key encoding and decoy state analysis.

The SNS protocol uses the Z basis for key bit encoding and the X basis for decoy state analysis. Conversely, the CAL protocol uses the X basis for encoding and the Z basis for decoy state analysis. This fundamental difference leads to important variations in practical implementation [16].

Decoy State Analysis

In the SNS protocol, the decoy state analysis is performed in the X basis via an interferometric measurement. For security reasons, the global phase is randomized by each terminal node at the beginning of each communication window. This requires Alice and Bob to reconcile their global phase after Charlie has communicated the outcome of his measurement. In the CAL protocol, on the other hand, the decoy state analysis is performed using the Z basis. Consequently, no global phase reconciliation is required between the two terminals, simplifying this part of the process.

Key Encoding (Signal Windows)

Regarding the encoding of key bits in the signal windows:

- In the **SNS protocol**, the Z basis is used for encoding. Alice and Bob's choice to send an empty state (with zero photons) or a state containing photons encodes bit 0 or 1, respectively.
- In the **CAL protocol**, the X basis is used for encoding. In this case, coherent states with two possible phases (that differ by π) are used, which, once interfered by Charlie, will cause a single photodetector to "click", to which bit 0 or 1 will be associated.

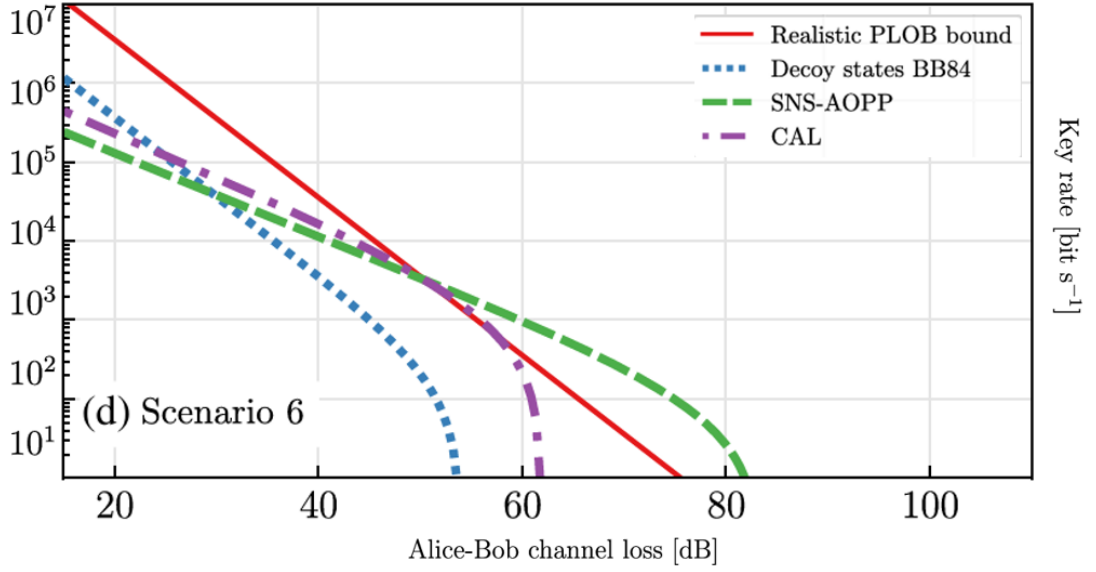


Figure 3.5: Simulated key rates for BB84, SNS-AOPP and CAL protocols as a function of channel loss, using a setup with ultra-stable independent lasers and SPAD detectors. This result is reported in ref.[16].

Achievable Key Rates and Protocol Performance

The notable difference in efficiency between the SNS-AOPP and CAL protocols, as illustrated in Fig. 3.5, stems from their distinct strategies for managing errors, which are essentially determined by the basis used for their signal and decoy windows. In particular, although both protocols are variants of TF-QKD, overcoming the PLOB bound (solid red line in Fig. 3.5) [8] and showing greater performance with respect to traditional QKD protocols like the BB84 with decoy states (dashed blue

line in Fig. 3.5), their performances diverge due to the basis (X or Z) used for decoy state analysis.

The CAL protocol performs its decoy state analysis in the Z basis, which is simpler to implement but less robust. In general, this approach limits the precision in estimating the Quantum Bit Error Rate (QBER). Consequently, the protocol's efficiency, and thus its Secure Key Rate (SKR), rapidly decreases as channel loss increases, limiting its effective range (dashed purple line in Fig. 3.5).

In contrast, the SNS protocol performs decoy state analysis in the X basis. While this measurement is technically more complex, it provides an exceptionally accurate estimation of phase errors and, consequently, a more robust SKR with respect to channel loss. This is a crucial property for applying TF-QKD in real-world telecommunication systems because phase fluctuations, as we will investigate in Chapter 4, are among the main sources of error in long-distance communications.

NOTE The addition of the AOPP (Active-Optics Phase-Partitioning) process further increases this robustness by efficiently identifying and discarding non-aligned bit pairs (as shown with a dashed green line in Fig. 3.5) [16]. This additional post-process filtering reduces the error rate and allows the protocol to maintain a high and stable SKR even over significantly longer distances.

3.5.3 Asymmetric SNS Protocol Implementation

In the context of quantum cryptography, the original version of the SNS protocol, as described by Ma et al. [34], presents a significant constraint in finite-key scenarios. To ensure security proof, it assumes that Alice and Bob use identical parameters for their sources, such as the intensities of the signal states ($s_A = s_B$) and the probabilities of sending coherent pulses in the Z windows ($\epsilon_A = \epsilon_B$). For this reason, this variant is often referred to as the *symmetric SNS protocol*.

To overcome this experimental constraint, which cannot always be satisfied in a practical implementation, an asymmetric version of the SNS protocol has been developed [28, 37]. This more recent and flexible variant does not require Alice and Bob to use identical source parameters, as shown in Figure [3.6]. However, this flexibility is conditioned on meeting specific additional conditions on the chosen parameters, which are indispensable for maintaining the validity of the security proof (see, for example, Equation 3.11 related to decoy states).

The adoption of the asymmetric version is not only motivated by the fact that it "relaxes" experimental requirements, but it has also been shown to offer a significantly higher key rate, particularly in scenarios with asymmetric communication channels. As highlighted in Figure [3.7], which shows the results obtained from different simulations reported in ref. [36], the key rate achieved with the asymmetric

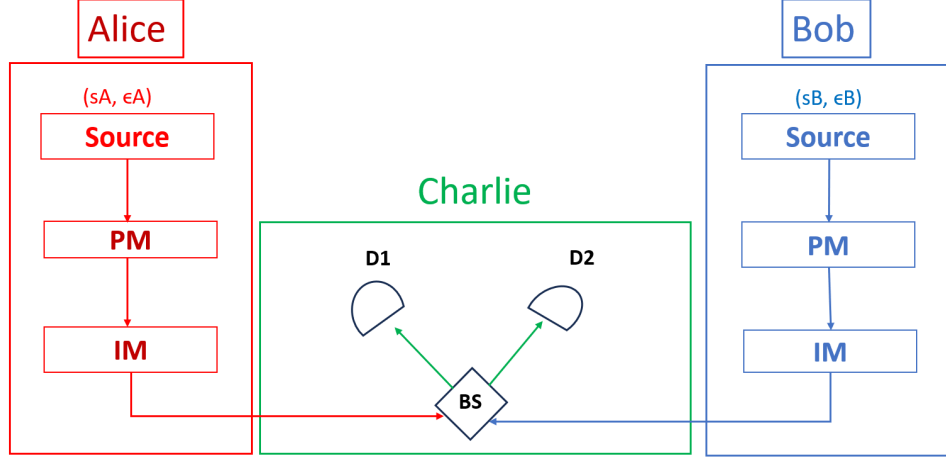


Figure 3.6: Schematic diagram of the setup for the SNS protocol. This diagram shows the setup for an SNS (*Send-No-Send*) protocol, where we indicate IM as intensity modulator; PM as phase modulator; BS as beam splitter; D1 & D2 as single-photon detectors in Charlie's measurement station. Unlike other protocols, in this diagram it is possible to implement an asymmetric version of the SNS, in which Alice and Bob are not constrained to use the same source parameters. They can therefore independently choose the intensity of the signal states (s_A, s_B) and the probability of sending coherent pulses in the Z windows (ϵ_A, ϵ_B), without them having to be equal.

protocol substantially surpasses that of the original symmetric version.

In this section, we will present the steps for the implementation of the most recent and efficient asymmetric variant of the SNS protocol. The original symmetric version can indeed be seen simply as a particular case of the asymmetric one, where Alice and Bob's source parameters are identical ($s_A = s_B$ and $\epsilon_A = \epsilon_B$).

NOTE For a deeper understanding of the foundations of the original protocol, one can refer to the work of Ma et al. [34].

Protocol Steps

Let's analyze the steps necessary for implementing this new asymmetric SNS TF-QKD protocol [28], including the decoy state method [38], and the mathematical constraints it imposes on the choice of the experimental parameters for our setup (Fig.[3.6]).

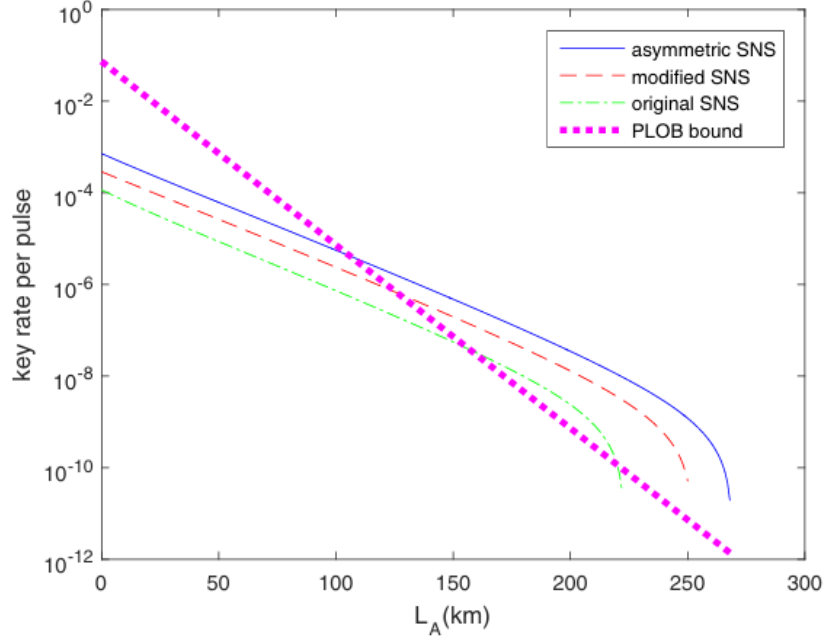


Figure 3.7: This graph shows the optimized key rates in relation to the transmission distance between Alice and Charlie obtained from different simulations reported in ref. [36]. The results of three different SNS protocols are compared: *asymmetric SNS*, *modified SNS*, and *original SNS*, while keeping the length difference between Alice and Bob’s channels fixed at 50 km. It is important to note that these results consider the effect of the finite size of the exchanged data. As shown, the performance of the asymmetric SNS protocol is significantly better than that of the original SNS, especially when the length difference between the channels is large. Furthermore, it can be observed that the asymmetric SNS protocol surpasses the linear theoretical limits for repeaterless QKD, such as the famous PLOB bound, demonstrating the significant advantage of the TF-QKD protocol over its original MDI-QKD version.

Step 1

In each i -th time window of the quantum communication, Alice and Bob independently decide, with a certain probability p_Z , to use the Z (signal) encoding basis. Alternatively, with the complementary probability $p_X = 1 - p_Z$, they choose the decoy X basis.

Decoy Window: For each decoy window, Alice (Bob) randomly chooses one of some previously defined states ρ_{A_k} (ρ_{B_k}), with $k = 0, 1, 2, \dots$

We define: $\rho_{A_0} = \rho_{B_0} = |0\rangle$ as vacuum states for $k = 0$ and ρ_{A_k} (ρ_{B_k}), for $k > 0$, as coherent states in the form [36]:

$$|\sqrt{\mu_{A_k}}e^{i(\delta_{A_i}+\gamma_{A_i})}\rangle, |\sqrt{\mu_{B_k}}e^{i(\delta_{B_i}+\gamma_{B_i})}\rangle \quad (3.8)$$

Where we denote with: μ_{A_k} (μ_{B_k}) the k -th intensity of a decoy state; with γ_{A_i} and γ_{B_i} the random phases and with δ_{A_i} (δ_{B_i}) the private phases.

Signal Window: For each signal window, Alice (Bob) can decide to:

- **Send a signal pulse** to Charlie with a probability ϵ_A (ϵ_B), which, analogous to the decoy ones, can be written in the state [36]:

$$|\sqrt{\mu'_A}e^{i(\delta_{A_i}+\gamma_{A_i})}\rangle, |\sqrt{\mu'_B}e^{i(\delta_{B_i}+\gamma_{B_i})}\rangle \quad (3.9)$$

Specifically, corresponding to this, she (he) decides to record a bit value of 1 (0). In this case, we denote with: μ'_A (μ'_B) the intensity of a signal state; with γ_{A_i} (γ_{B_i}) the random phase and δ_{A_i} (δ_{B_i}) the private phases.

- **Or decide not to send it** (or equivalently to send vacuum signal) :

$$\rho_{A_0}(\rho_{B_0}) = |0\rangle \quad (3.10)$$

Also in this case, she (he) records a correspondent bit value of 0 (1) with a probability of: $1 - \epsilon_A$ ($1 - \epsilon_B$).

There are a few crucial observations to highlight at this point.

- These recorded values, with the exception of the bits sacrificed during the post-processing phase, must remain secret for both parties. In fact, these bits will eventually form the final shared secret key.
- The probability of choosing the basis designated for the key, p_Z , must be maximized relative to p_X . This is essential for achieving a high key rate because decoy states are used exclusively in the X-basis for security proofs, specifically to detect potential beam splitter attacks. Similarly, to increase the probability of sending a single-photon pulse and mitigate the risk of attacks like photon number splitting, the probability of sending a vacuum state should be maximized. Transmitting a vacuum state (or a "non-send") minimizes the average intensity of the overall pulses, thereby enhancing transmission security.

- In the SNS protocol, the Z-basis is utilized for key encoding. In this basis, Alice and Bob transmit a *phase-randomized state* (Section 3.5.2). Crucially, the global phase of this state (e.g., $\gamma_i + \delta_i$) is never disclosed, which ensures the secrecy of the key bits. Conversely, the X-basis is reserved exclusively for decoy state analysis and security proofs. Users in this basis transmit a *phase-defined state* (Section 3.5.2). The global phase used for these states must be made public at the end of the quantum communication via the *Phase Reconciliation process* (Section 3.4.2), where Alice and Bob publicly announce their respective phase slide (Δ_k).

Definition 3.5.1 (Z, X, X(k) Windows). We define a **Z-window** when both Alice and Bob choose signal windows, and an **X-window** when both choose a decoy window. Among the X-windows, we define **X(k)-windows** as the slots in which both Alice and Bob choose the same decoy state, $\mu_{A,k}$ and $\mu_{B,k}$, where k is the same index that identifies the specific pair of decoy states used by them.

As can be seen from the two previous equations, for any type of pulse sent, they will always have to define an appropriate intensity and a relative phase. The intensities of the respective signal (decoy) states are previously agreed upon by Alice and Bob following an appropriate authentication of the parties and a classical communication.

Furthermore, as anticipated, for the specific implementation of the Asymmetric SNS Protocol, in which Alice and Bob can choose different source parameters such as the intensities of the signal states ($\mu'_A \neq \mu'_B$) and the probabilities of sending coherent pulses in the Z windows ($\epsilon_A \neq \epsilon_B$), the following mathematical constraint is required for the source parameters [36]:

$$\frac{\mu_{A_k}}{\mu_{B_k}} = \frac{\epsilon_A(1 - \epsilon_B)\mu'_A e^{-\mu'_A}}{\epsilon_B(1 - \epsilon_A)\mu'_B e^{-\mu'_B}} \quad \text{for each } k > 0. \quad (3.11)$$

NOTE This mathematical constraint is crucial for a secure implementation of the asymmetric SNS protocol and its purpose is in fact to ensure that the users' raw key is free from vulnerabilities such as systematic biases. In principle, it must be applied to all decoy pulses used.

Furthermore, note that this requirement is automatically satisfied for the symmetric protocol with: $\epsilon_A = \epsilon_B$ and $\mu'_A = \mu'_B$.

Discussion of Equation (3.11) Formally, the constraint in Eq. (3.11) ensures that the density matrix of the *untagged state* (not attacked by Eve) in the X-windows is the same as that of the untagged state in the Z-windows [39].

This condition is fundamental for using the decoy states technique, as it ensures

that the *bit-flip error rate* (Section 2.2.3) in the X-windows can be used to estimate the *phase-flip error rate* (Section 2.2.3) in the Z-windows [38]. Furthermore, it's important to note that, as anticipated, the sufficient condition for the security of this version of the SNS protocol is that all decoy pulses satisfy Eq. (3.11).

However, in this specific work, which uses the four-intensity decoy state method [36], not all μ_{A_k} and μ_{B_k} are required to satisfy this equation. In fact, only the decoy states μ_{A_1} and μ_{B_1} will be used to estimate the phase-flip error rate (e_{ph}) in the X-windows, and consequently, only for these two must we ensure compliance with Eq. (3.11) [38].

Step 2: Charlie's Measurement and Public Announcement

Charlie, acting as the central measurement station, executes an interferometric measurement on the twin quantum fields transmitted by Alice and Bob during each dedicated communication window. Specifically, he records the clicks registered on his single-photon detectors (D_0 and D_1) and then, following the quantum transmission phase, he publicly a data string containing the complete record of these detector clicks across all time windows utilized for the QKD protocol.

NOTE The click announced by Charlie can result from three main scenarios, each with a different impact on key generation:

- **Constructive Interference:** The click is the result of interference between signals sent simultaneously by Alice and Bob. Although this is the ideal constructive interference signal, its probability is very small, due to the high probability that either Alice or Bob chose to send a vacuum state.
- **Single Transmission:** The click results from a signal sent by only one participant (Alice or Bob, but not both). This is the most probable event that, once reconciled, contributes directly to the raw key.
- **Background Noise:** The click occurs when neither participant sends a signal, and is due only to the intrinsic detector noise (dark count). This scenario does not contribute to the key and introduces a potential error that will have to be managed in the subsequent phases of the protocol.

Definition 3.5.2 (Effective Window). An **Effective Event** is defined as an event in which one and only one of Charlie's detectors produces a click. Consequently, the corresponding time window is called an **Effective Window**.

Only the data collected from these Effective Windows will be used by Alice and Bob for:

- Key distillation from the Z Effective Window [8].
- Estimation of key channel parameters, such as the number of untagged bits (n_1) and the phase-flip error (e_{ph}), from the X Effective Window [38].

Consequently, Alice and Bob will consider only the events announced by a single detector. They will select from Charlie's data only the windows in which a single event was recorded, as only these can contribute to the final key or the decoy state analysis. On the contrary, events with double clicks or zero clicks are discarded, contributing to the QBER caused by bit-flips (Section 2.2.3).

After Alice and Bob have repeated the previous steps N times to accumulate sufficient data, the secret key distillation process begins.

In particular, from the string of n_t bits derived from the Z Effective Windows, Alice and Bob will perform the following data post-processing phases: **Sifting, Parameter Estimation, Error Correction, and Privacy Amplification**.

Step 3: Sifting

In this step, Alice and Bob publicly announce:

- Which windows were chosen as Signal and Decoy.
- For Decoy windows, they also announce the intensities and the respective private phases, δ_A (δ_B), that were chosen for each decoy pulse.

Property: Photon Number Distribution of Phase-Randomized States

When Alice and Bob choose the Z basis for transmission, they send a phase-randomized coherent state (described in Section 3.5.2). In this basis, the coherent state transmitted by Alice (with intensity μ'_A) or Bob (with intensity μ'_B) can be mathematically expressed as a density matrix in the photon number space[36]:

$$\rho = \sum_{k=0}^{\infty} \frac{e^{-\mu'} \mu'^k}{k!} |k\rangle\langle k| \quad (3.12)$$

where μ' represents the average intensity of the sent signal (which will be μ'_A or μ'_B depending on the sender). This means that the state appears as a classical

statistical mixture of states with a variable photon number k , following a Poisson distribution.

Since the security of the QKD protocol is guaranteed only by single-photon contributions ($k = 1$), and given the emission of a mixture of states with different k , Alice and Bob define a restricted subset of all Z windows where a click was registered. This subset is called a **Z1-Window** and corresponds to the *untagged* events from which the key will be distilled.

Definition 3.5.3 (Z1-Window). Z1-Windows are defined as all the windows in which:

- **One and only one** of the two parties (Alice or Bob) sends a pulse.
- The pulse sent **is a single-photon state**.

Unlike the Z Effective Windows, which are identified from Charlie's public announcement, Alice and Bob do not know which time window is a Z1-window (this is due to the inherent quantum nature of their sources, as described by the Poisson distribution in Eq. (2.1)).

However, they can calculate the number of Z1-windows using the Decoy State technique [38]. This allows them to estimate the lower bound of the number of untagged bits, which originate from pulses where only one party actually sent a single photon. Conversely, all multi-photon signals must be considered *tagged*, as they are vulnerable to a beam splitter attack by Eve.

Definition 3.5.4 (Z1-Effective-Window). The *Z1-Effective-Window* is the final subset that combines the *Z1-Window* with the *Z-Effective-Window*. A Z1-Effective Window thus corresponds to an event in which:

- **one and only one** of Alice and Bob has sent a single-photon state;
- and **only one** of Charlie's detectors has clicked.

It is from this specific subset that the bits of the shared key will be derived, which are, in principle, identical and completely secret.

Definition 3.5.5 (X(k)-Effective-Window). Within the set of X-windows, a subset of *X(k)-Effective-Windows* is defined as those in which:

- The intensities μ_{A_k} and μ_{B_k} with the same index k are chosen.

- The phase shifts δ_A and δ_B satisfy the restriction [8]:

$$|\delta_A - \delta_B| \leq \frac{2\pi}{M} \quad \text{or} \quad |\delta_A - \delta_B - \pi| \leq \frac{2\pi}{M} \quad (3.13)$$

Note on the Second Condition. One of the most important characteristics of the TF-QKD protocol is that, although it uses the same components as the well-known MDI-QKD protocol, it also requires the coordinated randomization of the twin-field phases (3.4.2). Specifically, Eq. (3.13) is what allows Alice and Bob to reconcile their two private phases, which are chosen randomly and independently from each other for each quantum transmission window within the semi-open interval $[0, 2\pi)$.

Step 4: Parameter Estimation

At this point, Alice and Bob process the data collected from the effective Z and X windows to derive some fundamental quantities for the calculation of the final Secret Key Rate (SKR).

First, they randomly choose and publicly declare a small sample of effective Z windows from the string of n_t bits announced by Charlie. It is important to underline that this subset of bits must be discarded, as it is now known to any potential eavesdropper, but it is used to estimate the bit-flip error rate, E_Z .

A **bit-flip error** occurs when Alice's bit value differs from Bob's in a Z window. The bit-flip error rate E_Z can therefore be written as [8]:

$$E_Z = \frac{n_{NN} + n_{SS}}{n_t}$$

Where:

- n_t is the length of the string announced by Charlie from the events in the effective Z windows.
- n_{SS} is the number of times they both sent.
- n_{NN} is the number of times they both did not send.

Based on the data collected from the effective X(k)-Effective-Windows—that is, on the announced values of δ_A and δ_B for these windows—they estimate two quantities crucial for the decoy-state method: n_1 and e_{ph1} [38].

Definition 3.5.6. n_1 is defined as the number of events in the effective Z1 windows, which corresponds to the number of untagged bits used to generate the final key.

Definition 3.5.7. e_{ph1} is defined as the single-photon phase-flip error rate of the states in the effective Z1 windows.

NOTE. This is a fundamental step that relies on the use of the **Four-intensity decoy-state analysis**. This method requires observing the counting rates of various input light intensities in the effective X windows to estimate the lower bound of n_1 . Conversely, only the measurement results from the effective X(k=1) windows are used to estimate the lower bound of e_{ph1} [36].

Finally, for the calculation of n_1 and e_{ph1} , it is important to remember that:

- **In the asymptotic case** where decoy states with an infinite number of different intensities are used, it is possible to obtain the **exact value** of e_{ph1} and n_1 .
- **In the real-world case**, with a finite number of decoy intensities, the *finite key size effect* means it is only possible to obtain their respective **lower and upper bounds**. From these bounds, an estimate of the real values of n_1 and e_{ph1} is then derived [36].

Step 5: Error Correction

At this stage, Alice and Bob have formed two bit strings, Z_s and Z'_s , each of length n_t , based on events in the Z set.

While these two strings would be identical in an ideal world, the inherent noise in a real communication channel means there is a non-zero probability that they will differ. Alice, therefore, needs to perform an information reconciliation scheme to correct Z'_s . Specifically, the goal is to make the raw bit strings, which contain errors from the quantum channel, identical for both Alice (Z_s) and Bob (Z'_s) while minimizing information leakage to a potential eavesdropper, Eve.

The essential idea to perform this error correction is as follows [40]:

1. Alice initiates the scheme using a specific error correction protocol. She calculates a minimal set of correction data based on her string Z_s . This data, known as **leak_EC**, is then made public and accessible to both Bob and any potential eavesdropper.

2. Next, Alice randomly selects a universal hash function and applies it to her string Z_s to derive a hash value. The length of this hash is defined by the required security level, specified by the term $\log_2(1/\epsilon_{cor})$.
3. Alice transmits both the chosen hash function and the calculated hash value to Bob.
4. Bob uses the public correction bits `leak_EC` to rectify the errors in his initial sequence Z'_s , thereby producing his new rectified string, named $\hat{Z}_s$.
5. Bob applies the same universal hash function, received from Alice, to his corrected string, $\hat{Z}_s$.
6. Finally, Bob compares his calculated hash with the one transmitted to him by Alice. If the two hash values match, reconciliation is considered a success. In this scenario, Alice and Bob assume that their final strings (Z_s and $\hat{Z}_s$) are consistent, with a residual error probability less than ϵ_{cor} . Conversely, if the values do not match, this indicates an excessive error rate or a potential attack, rendering the correction reliability null.

NOTE The existence of a non-zero failure probability, ϵ_{cor} , even after successful reconciliation, is a consequence of the *finite-key size effect* inherent in the asymmetric SNS protocol. This effect arises because real-world QKD systems operate with a limited amount of data, unlike the idealized infinite-data scenarios [8].

Step 6: Privacy Amplification

The final phase of the protocol is Privacy Amplification (PA), which aims to produce a key that is unconditionally secure (Section 1.2.3). Following successful error correction (Step 5), Alice and Bob share an identical but partially compromised string, α , that can be considered as the provisional key.

The public communication required for error correction inevitably provided an eavesdropper (Eve) with some residual information about α . Consequently, the primary goal of PA is to eliminate this information leakage and ensure Eve's total ignorance of the final secret key. This is achieved by extracting a shorter, compressed string, α' , of length l , from the longer string α . Only this shorter string will constitute the shared secret key.

The standard technique for PA utilizes universal families of hash functions. The procedure involves the following steps [2]:

1. Alice and Bob publicly and randomly select a universal hash functions, h .

2. Both parties apply this function to the provisional key, calculating the final key as $\alpha' = h(\alpha)$.

NOTE The security of the final key is directly proportional to the number of sacrificed bits, thereby establishing a clear trade-off between the resulting key length and its level of unconditional secrecy.

Step 7: Estimation of the Key Rate per Time Window with the Effect of Finite Key Size

It is possible now to present the formula that Alice and Bob can use to calculate the key rate considering the effect of the finite key size in the case of our described asymmetric SNS protocol.

NOTE While we have limited our discussion so far to the asymptotic case, where the size of the exchanged data is infinite and observed values match their expected counterparts, a real QKD experiment operates in a non-asymptotic scenario. In this case, the exchanged data has a finite size, and the observed values of parameters derived from decoy-state analysis will differ from their expected values. To bridge this gap between observed and expected values, we will introduce the **Chernoff bound** [28] that allows us to estimate the expected values from the observed ones and then use a worst-case scenario analysis to ensure the final key is secure.

The length of the final key, N_f , to be secure, must satisfy the following condition [36]:

$$N_f = n_1 [1 - H(e_{ph1})] - f n_t H(E_Z) - \log_2 \frac{2}{\epsilon_{cor}} - 2 \log_2 \sqrt{\frac{1}{2\epsilon_P A \hat{\epsilon}}} \quad (3.14)$$

Where:

- f is the inefficiency of the correction.
- n_t is the number of effective Z windows.
- E_Z is the bit-flip error rate in the effective Z windows.
- ϵ_{cor} is the probability that error correction fails.
- ϵ^- is the probability that the real value of e_{ph1} does not fall within the estimated interval.
- ϵ_{PA} is the failure probability of privacy amplification.
- ϵ_{n1} is the probability that the real value of n_1 does not fall within the estimated interval.

The presented Eq.(3.14) that determines the final secure key length (N_f) is based on a critical balance between the acquired information and the penalties imposed

by practical security procedures. The main term, $(n_1[1 - H(e_{ph1})])$, quantifies the maximum amount of secret that can be extracted. It depends on the estimation of effective single-photon events (n_1) and Eve's residual uncertainty, defined by the quantum error rate ($H(e_{ph1})$).

From this gain, two essential penalties must be subtracted. The first is the error correction loss ($f n_t H(E_Z)$) (Step 5) and the second is the privacy amplification penalty ($2 \log_2(1/\epsilon_{cor})$) (Step 6).

Finally, there is the finite-key penalty ($2 \log_2 \sqrt{\frac{1}{2\epsilon_{PA}\hat{\epsilon}}}$): that reflects the statistical uncertainty introduced by the fact that Alice and Bob use only a finite sample of data to estimate the critical security parameters (n_1 and e_{ph1}) [36].

Equation (3.14) can be written in the form of key rate per time window with some source parameters [36]:

$$\begin{aligned}
 R = p_A^Z p_B^Z \bigg\{ & [\epsilon_A(1 - \epsilon_B)\mu'_A e^{-\mu'_A} + \epsilon_B(1 - \epsilon_A)\mu'_B e^{-\mu'_B}] \\
 & \times s_1^Z [1 - H(e_1^{ph})] - f S_Z H(E_Z) \\
 & - \frac{1}{N_t} \left(\log_2 \frac{2}{\epsilon_{cor}} + 2 \log_2 \frac{1}{\sqrt{2\epsilon_{PA}\hat{\epsilon}}} \right) \bigg\}
 \end{aligned} \tag{3.15}$$

NOTE. Once again, if we were to choose for Alice and Bob the particular case where $\mu'_A = \mu'_B$ and $\epsilon_A = \epsilon_B$, we would return to the original symmetric SNS protocol.

Chapter 4

TF-QKD: Challenges and Real-World Implementations

4.1 The Problem of Phase Coherence

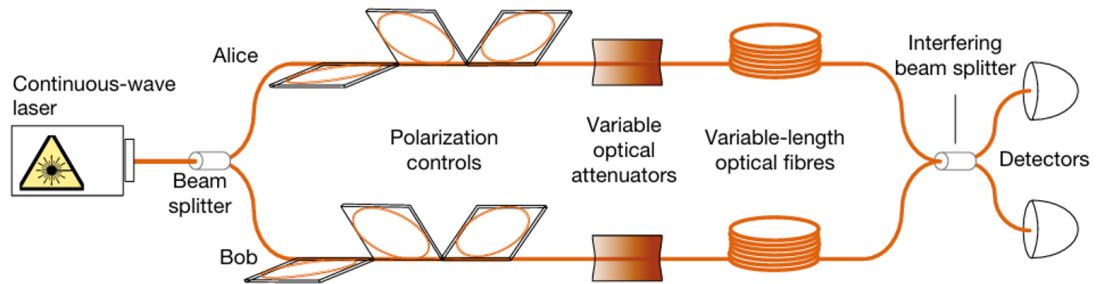


Figure 4.1: A schematic of the experimental setup used in [29] to measure phase drift in a Mach-Zehnder interferometer with long optical fibers. A continuous-wave laser is split, and the two pulses travel along separate paths to an interfering beam splitter. The setup includes polarization controls, variable optical attenuators, and variable-length optical fibers to simulate different conditions.

4.1.1 TF-QKD Critical Assumptions

The primary technical obstacle in the real-world implementation of the Twin-Field Quantum Key Distribution (TF-QKD) protocol lies in the extremely precise control of the relative phase evolution of the twin optical fields. To ensure an effective real-world implementation, two fundamental critical assumptions must be met [12]:

1. The twin fields must be phase-coherent at the moment of their generation by Alice and Bob.

2. They must preserve this coherence along the entire path up to Charlie's detector.

To satisfy these two assumptions, it is fundamental to study the evolution of the relative phase of the twin fields generated by Alice and Bob. In particular, we need to take into account that this relative phase inevitably undergoes fluctuations induced by environmental noise when the fields travel hundreds of kilometers over the two distinct optical fiber paths placed between the terminal node and the central one.

4.1.2 Sources of Phase Fluctuation

The differential phase evolution $\Delta\phi(t)$ between the two optical paths can be classically expressed as a combination of static and dynamic imbalances. Generally, it can be modeled by two main terms [16]:

$$\Delta\phi(t) = 2\pi(\Delta\nu \cdot t + \nu \cdot \Delta L/c) \quad (4.1)$$

where c is the speed of light in vacuum; ν is the nominal optical frequency of the fields ($\nu_A \approx \nu_B \approx \nu$); $\Delta\nu = \nu_A - \nu_B$ is the static frequency imbalance between sources A and B; and ΔL is the static length imbalance between the two optical fiber paths. Let's analyze the individual contributions.

1. Static Frequency Imbalance ($2\pi\Delta\nu \cdot t$) This term arises from the static frequency imbalance ($\Delta\nu = \nu_A - \nu_B$) between Alice's and Bob's laser sources, which can rapidly destroy the coherence condition. This problem is effectively mitigated today through the use of Phase-Locked Loops (see section 4.3) (PLLs) [41]. The goal is to stabilize the two lasers until $\Delta\nu < 1$ Hz is achieved, a value that ensures a negligible contribution to the Quantum Bit Error Rate (QBER).

2. Dynamic Optical Path Fluctuation ($2\pi\nu \cdot \Delta L/c$) The second and most critical problem is the dynamic fluctuation of the relative phase ($\delta\phi(t)$) induced by environmental noise.

To define this noise, it is necessary to introduce the concept of Optical Path Length (OPL) [15]. The OPL is the effective physical length of the optical path and can be defined as the product of the refractive index of the medium (n) and its physical length (L):

$$\text{OPL} = nL$$

The phase noise $\delta\phi(t)$ is directly proportional to the dynamic fluctuation of the OPL, $\delta(nL)$. The total fluctuation $\delta(nL)$ comprises two fundamental contributions

[42]:

$$\delta(nL) = n \cdot \delta L + L \cdot \delta n$$

- $n \cdot \delta L$: This term is related to fluctuations in the physical length (δL) of the optical fiber, often caused by mechanical stress or thermal noise.
- $L \cdot \delta n$: This term is related to fluctuations in the refractive index of the optical fiber (δn), caused by variations in ambient temperature and pressure.

While the term describing the static frequency imbalance is treated as a solvable problem (via PLLs), the element representing the true technological challenge is the dynamic fluctuation of the OPL, $\delta(nL)$. In fact, this temporal variation constitutes the dominant noise source, thus requiring continuous and active compensation for long-range TF-QKD systems to operate successfully.

The magnitude of this issue is confirmed by experimental data. In setups with relatively short fiber arms, such as those employing 40 km of fiber (e.g., a Mach-Zehnder interferometer Fig. [4.1]), the phase drift rate is manageable, registering approximately $0.1 \text{ rad} \cdot \text{ms}^{-1}$ under uncompensated conditions (as reported in ref.[39]). However, extending the distance, the noise scales exponentially: the results reported in ref. [39] show that the drift rate increases significantly, reaching $2.4 \text{ rad} \cdot \text{ms}^{-1}$ for a total distance of 100 km and peaking at $6.0 \text{ rad} \cdot \text{ms}^{-1}$ over a critical distance of 550 km.

This showed dependence on distance unequivocally confirms that the TF-QKD protocol's performance is directly deteriorated by environmental perturbations acting independently on the two optical paths. In particular, the increase in phase drift velocity with longer links highlights the impossibility of implementing the TF-QKD protocol over long distances without adopting a robust system dedicated to phase realignment and stabilization.

4.1.3 Impact of Phase Stability on QBER

Following this analysis, it is clear that, establishing high phase coherence between the twin fields emitted by Alice and Bob is an indispensable requirement for the effectiveness of the TF-QKD protocol. The success of the protocol relies on the ability to achieve perfect constructive or destructive interference at Charlie's detector. Any instability in the relative phase directly translates into an increase in the QBER and, consequently, a reduction in the achievable SKR.

To quantify the impact of this dynamic imperfection, the variance of the phase fluctuations (σ_ϕ^2) observed at Charlie's detector is considered. The specific contribution of this phase noise (e_ϕ) to the total QBER is modeled by the following

relationship [16]:

$$e_\phi = \int \sin^2(2\phi)P(\phi)d\phi \approx 4\sigma_\phi^2 \quad (4.2)$$

NOTE This linear approximation is valid only by assuming that the phase fluctuations follow a Gaussian distribution [16].

4.2 Real-World Phase Stabilization for TF-QKD

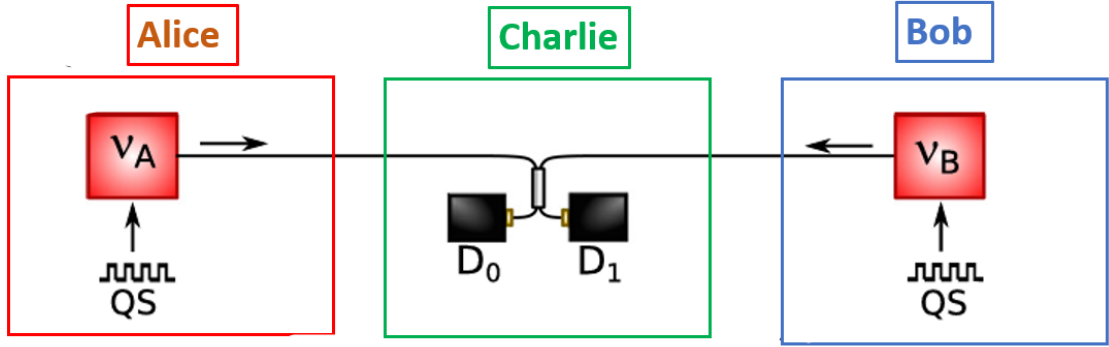


Figure 4.2: Ideal Scheme of the TF-QKD Protocol [16]. Alice and Bob generate quantum states (QS) using two local lasers, attenuated to a single-photon level via variable optical attenuators (VOA). They modulate their respective quantum signals using intensity and phase modulators based on the specific TF-QKD protocol variant chosen (Section 3.3). In this ideal scenario, Alice’s and Bob’s quantum lasers have perfectly equal frequencies ($\nu_A = \nu_B$) and maintain total phase coherence. The signals are sent via two quantum channels of equal length (L) toward Charlie’s central station, where they are measured by single-photon detectors (D_0 and D_1). This theoretical scheme automatically satisfies the two critical real-world TF-QKD requirements (Section 4.1.1).

We will now explore some practical architectures developed to maintain phase coherence, focusing on the most common topologies for a real-world implementation of the TF-QKD protocol.

To define a reference point for these practical solutions, we first present an ideal scheme of the protocol in Fig.[4.2]. In this hypothetical scenario, the two critical assumptions for TF-QKD are automatically satisfied in the absence of intrinsic laser noise or environmental disturbance. Specifically:

1. **Perfect Coherence:** Alice’s and Bob’s quantum lasers are characterized by perfect, continuous phase coherence, ensuring their frequencies are perfectly matched ($\nu_A = \nu_B$) at the central station.

2. Perfect Symmetry: The resulting quantum signals are sent via two perfectly symmetric quantum channels of equal length (L) to Charlie's central station, eliminating any relative phase drift induced by path length differences.

This idealized configuration allows for maximum secret key generation and can be used as the theoretical upper limit for the following real-world TF-QKD implementations.

4.2.1 Spectral Phase Analysis

Before we proceed by investigating the most common topologies for a real-world TF-QKD implementation, it's essential to precisely define the relationship between the phase variance (σ_ϕ^2), the QBER contribution (e_ϕ), and the temporal interval for key generation (τ_Q). To mathematically quantify the connection between σ_ϕ and τ_Q , we follow the demonstration reported in ref.[16]. For any given variable $y(t)$, this is defined as:

$$S_y(f) = \mathcal{F}[R(y)] \quad (4.3)$$

Where $\mathcal{F}$ represents the Fourier transform, and $R(y)$ is the autocorrelation function of the variable $y(t)$. This quantity is particularly useful because, according to the Wiener-Khintchine theorem [16], the phase variance σ_ϕ can be expressed directly in terms of the phase noise power spectral density $S_\phi(f)$:

$$\sigma_\phi^2(\tau_Q) = \langle \Delta^2 \phi \rangle_{\tau_Q} = \int_{1/\tau_Q}^{\infty} S_\phi(f) df \quad (4.4)$$

Equation (4.4) [16] is fundamental to our work (Chapter 7), as it allows us to derive σ_ϕ in various TF-QKD topologies by measuring $S_\phi(f)$.

As we have previously noted, the total $S_\phi(f)$ is predominantly determined by two main contributions: the phase fluctuations accumulated by photons traveling through telecommunication fibers due to environmental noise and the initial phase misalignment between the twin photons generated by Alice and Bob's distinct laser sources. These two contributions will ultimately determine the efficiency of the phase correction schemes we will introduce in the next section.

4.2.2 Common-Laser Scheme

Introduction

The first scheme proposed to address the two critical assumptions of TF-QKD (Section 4.1.1) is the Common-Laser Scheme, illustrated in Figure [4.3]. The central idea of the Common-Laser Scheme is to ensure phase coherence between

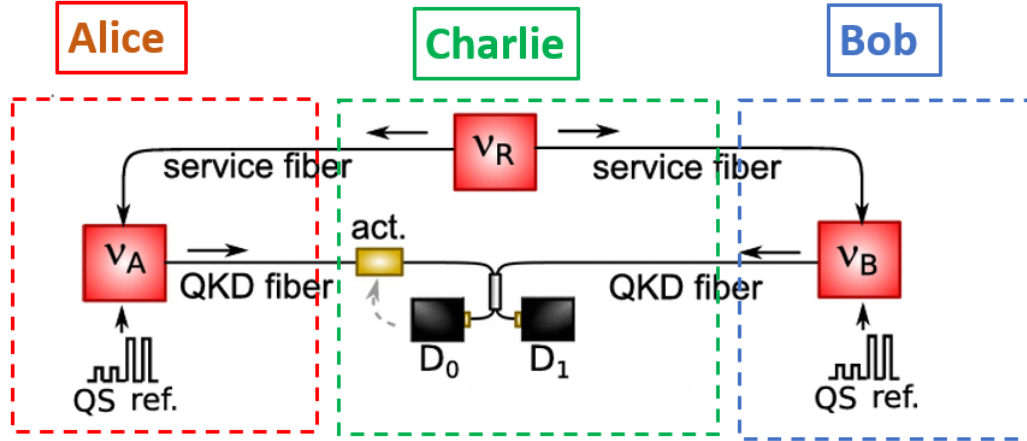


Figure 4.3: Common-Laser Scheme for the TF-QKD protocol [16]. In this scheme, a reference laser (*Ref. laser*) positioned at the central station of Charlie provides a stable frequency (ν_R) to Alice and Bob through a dedicated *service channel*. This ensures that Alice’s and Bob’s local lasers have identical frequencies ($\nu_A = \nu_B = \nu_R$), satisfying the first critical assumption of TF-QKD. Furthermore, to address dynamic phase noise introduced by environmental variations (e.g., temperature, vibrations), the protocol employs an interleaving mechanism [43]. This periodically interrupts QKD (QS pulses) to switch to a classical regime where Charlie sends more intense reference pulses (Ref. pulses). By measuring the phase shift of these pulses, Charlie can accurately determine the relative phase error between Alice’s and Bob’s twin fields and apply the necessary real-time correction.

the twin fields at the source (satisfying the first critical assumption of TF-QKD) by using a reference laser (ν_R), typically located at Charlie. This laser is phase-locked to the two local lasers of Alice and Bob via a dedicated service channel. In this way, both peripheral sources effectively "copy" the frequency of the reference laser ($\nu_A = \nu_B = \nu_R$), eliminating the static frequency imbalance $\Delta\nu$. However, satisfying the second critical assumption of TF-QKD — i.e., the preservation of phase coherence towards the central node — remains the main obstacle.

Detailed Analysis of Residual Phase Noise Sources

Despite the Common-Laser approach typically positioning the reference laser (ν_R) at Charlie to mitigate degradation due to laser frequency fluctuations, perfect phase stability is never achieved. Even in an ideal condition where Alice’s (ν_A) and Bob’s (ν_B) quantum lasers were to perfectly copy the frequency of the distributed reference laser (ν_R), dynamically induced phase fluctuations would persist along the

two optical fibers. Interference at Charlie's central node therefore occurs between two signals of identical frequency, but which have accumulated uncorrelated phase noise contributions along their respective paths. Consequently, the interfering photons will always show residual phase fluctuations, whose spectral power density, $S_\phi(f)$, is a crucial element for the design of the TF-QKD protocol.

More precisely, the spectral power density of the residual phase noise, $S_\phi(f)$, is the sum of two main contributions [16]:

$$S_\phi(f) = 4 \sin^2 \left(\frac{2\pi f n \Delta L}{c} \right) S_{l,C}(f) + 4[S_{F,A}(f) + S_{F,B}(f)] \quad (4.5)$$

Where f is the Fourier frequency (Hz); n is the effective refractive index of the optical fiber; ΔL is the geometric length imbalance between the two fiber arms ($L_{AC} - L_{BC}$) (m); c is the speed of light in a vacuum (m/s); $S_{l,C}(f)$ is the spectral power density of the distributed reference laser noise (rad^2/Hz); and finally $S_{F,A}(f)$, $S_{F,B}(f)$ are the spectral power densities of the dynamic phase noise induced by the optical fibers along the Alice-Charlie and Bob-Charlie paths (rad^2/Hz).

The two terms represent distinct noise sources:

- **Laser Noise:** The first term, $S_{l,C}(f)$, derives from the self-delayed interference of the reference laser itself [16]. This noise emerges due to an imperfect geometric identity of the optical channels (an imbalance, ΔL) connecting Alice and Bob to Charlie. Consequently, this noise term is zero under perfect balancing conditions ($\Delta L = 0$), but increases with the increment of length deviation.

Furthermore, this term is strictly dependent on the coherence of the reference laser. For the stringent requirements of TF-QKD, it is therefore common practice to use ultra-stable high-coherence lasers.

- **Fiber Noise:** The second term of Equation (4.5) represents the noise originating from the optical fibers that connect Alice and Bob to Charlie. This noise is highly dependent on the surrounding environment (for example, metropolitan areas with heavy traffic show higher noise levels) and is, in particular, often the dominant term in Equation (4.5), especially over long distances.

The corresponding spectral power density of this phase noise, $S_F(f, L)$, is approximately given by [16]:

$$S_F(f, L) = lL \frac{f^2}{(f + f'_c)^2} \quad (4.6)$$

The most important characteristic of this noise is its linear scaling with the fiber length (L), mediated by an empirical coefficient (l) that depends on the

fiber type and the environment. At low frequencies ($f \ll f'_c$), the noise scales as f^2 , while at high frequencies ($f \gg f'_c$) it stabilizes at a constant value $l \cdot L$ [16].

Dynamic Phase Stabilization via Interleaved Technique

The correction of uncorrelated dynamic fluctuations in the length and refractive index of the two optical fibers, caused by environmental noise, is essential to maintain a high interference visibility at Charlie's node. To overcome this problem, proof-of-principle experiments [43] use an interleaving technique that alternates quantum transmission with classical transmission. The purpose of this classical phase is to measure the relative phase error between Alice's and Bob's Twin Field pulses, thus allowing Charlie to determine it accurately and apply the necessary correction in real time. The stabilization approach is based on the alternation of

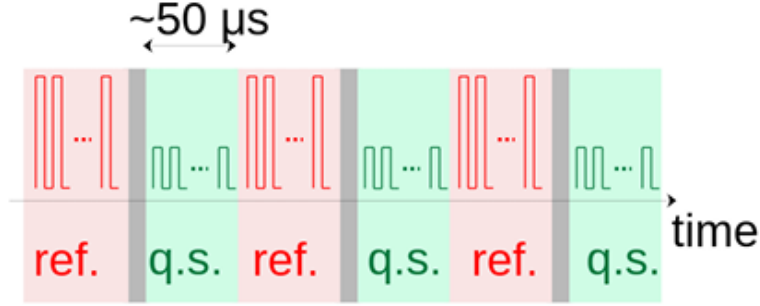


Figure 4.4: Time Scheme for Classical-Quantum Interleaving in TF-QKD [12]. Due to the rapid increase in relative phase noise and the consequent degradation of the QBER, the unstabilized TF-QKD protocol cannot operate continuously. Communication is divided into cycles: short quantum windows (q.s., quantum state), typically limited to $\sim 50 \mu\text{s}$ or less [43], are dedicated to sending the photonic signals for key generation. These windows are interleaved with classical reference periods (ref., reference), during which more intense (classical) light pulses are sent for monitoring and active realignment of the relative phase $\Delta\phi$ between Alice and Bob. This need for frequent realignment drastically limits the useful time for key extraction, reducing the overall efficiency of the SKR.

two distinct operating regimes that define a duty cycle:

1. **Quantum Regime (τ_Q):** For a time interval τ_Q , during which the phase is sufficiently stable, the system operates in the quantum regime. The laser source (the quantum laser) and the modulators are used to prepare and send low-photon-number states (for example, Signal and Decoy states in the SNS

protocol, as described in Section 3.5.3), used for the potential exchange of key bits or to detect the presence of an eavesdropper. For this purpose, the VOA (Variable Optical Attenuator) (Fig.[3.3]) applies a strong attenuation to achieve the required quantum regime.

2. **Classical/Sensing Regime (τ_{PS}):** When the relative phase is no longer stable and realignment is necessary, the QKD protocol is interrupted and the system switches to a classical regime. This transition is enabled by reducing the laser attenuation via the VOA, effectively transforming the quantum laser into a sensing laser. Charlie uses this sensing laser to detect the relative phase fluctuation between the two arms of the interferometer. Based on this information, he calculates the necessary correction and applies it via a phase modulator placed on one of the arms, thus compensating for the relative drift between the twin fields.

The most significant limitation of this solution is the introduction of a dead time (τ_{PS}) required for the phase realignment procedure, which inevitably reduces the secure key rate. Mathematically, this effect is quantified by the duty cycle (d), which must be multiplied by the theoretical secret key rate achieved by the considered protocol variant to determine the effective key rate [16].

$$d = \frac{\tau_Q}{\tau_Q + \tau_{PS}} \quad (4.7)$$

Where τ_Q is the maximum time of uninterrupted QKD transmission and τ_{PS} is the time required for phase stabilization. The quantum integration time (τ_Q) is inversely correlated with the residual phase variance (σ_ϕ^2): reducing the system's phase noise allows for an increase in τ_Q , which directly increases the Secret Key Rate (SKR). In the absence of active and effective phase stabilization, environmental and intrinsic fluctuations induce a very rapid increase in the Quantum Bit Error Rate (QBER). Experimental results reported in Ref. [41] show that, in an unstabilized scenario, the system exceeds the critical QBER threshold of 1% in approximately 100 μ s. Consequently, to maintain signal integrity, the integration time on Charlie's detectors must be limited to very short quantum windows (τ_Q), typically on the order of 20-50 μ s [43]. This drastic limitation on the duration of τ_Q has a significant impact, reducing the achievable SKR, as a large portion of the total time is spent on compensating for the relative phase instability of the TF-QKD system [12].

Limitations for Long-Distance Transmissions

Although the Common-Laser Scheme is a fundamental solution for implementing the TF-QKD protocol, its effectiveness drastically decreases with increasing communication distance [16]. The main cause lies in the second term of Equation (4.5)

— the fiber noise — which scales linearly with length, as highlighted by Equation (4.6). For distances exceeding a few hundred kilometers, where attenuation and phase fluctuations are significantly higher, the interleaving approach is no longer able to suppress dynamic noise fast enough to keep the QBER below the security threshold.

In conclusion, the Common-Laser Scheme effectively solves only the first of the two critical assumptions (coherence at the source). For long distances, it is insufficient to optimally manage the dominant phase noise generated by signal propagation in the fibers. This leads to a drastic reduction in the QKD transmission duty cycle and, consequently, in the achievable secure key rate [41].

4.2.3 Common Laser Scheme with Dual-Band Stabilization

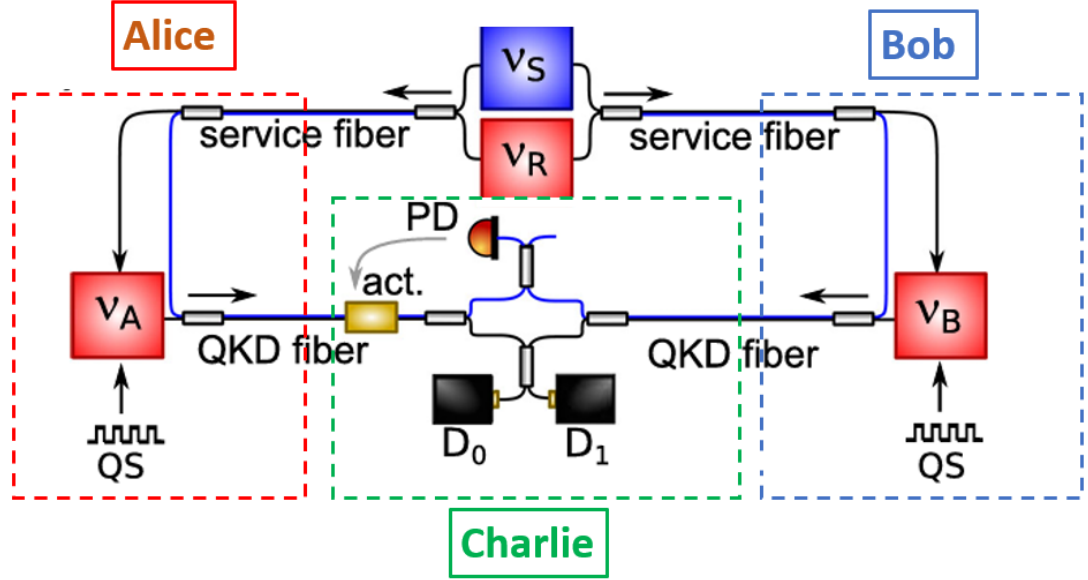


Figure 4.5: The Common Laser scheme with dual-band stabilization for TF-QKD. [44] A reference laser (ν_R) and a sensing laser (ν_S) are used to stabilize the phases of Alice's (ν_A) and Bob's (ν_B) local lasers, which are used for quantum state (QS) preparation. The setup employs two distinct fiber types: *QKD fibers* for quantum signal transmission and *service fibers* for laser synchronization. Charlie's station includes a noise detection and cancellation (NDC) system, a photodiode (PD) for phase monitoring, and an actuator (act.) for active fiber stabilization, ensuring high secret key rates over long distances.

Introduction to Dual-Band Stabilization

Having introduced the operating principle and the main limitations of the Common-Laser Scheme (CLS), we now present an evolution of that scheme in Figure [4.5] [41].

This new approach employs *Dual-Band Stabilization*. By leveraging spectral separation methods, such as classical Dense Wavelength Division Multiplexing (DWDM), it enables active phase noise cancellation. The main advantage of this innovative methodology is the ability to significantly reduce the need for periodic realignments (or dead time τ_{PS}), allowing the system to operate with a very high duty cycle ($d > 0.9$) [16].

The fundamental principle of the scheme lies in the addition of an auxiliary *sensing laser* to the Common-Laser setup (Fig. [4.5]). Operating at a slightly different wavelength (ν_S), this laser is sent along the same optical fiber as the quantum laser using a DWDM filter. The sensing laser is specifically designed to detect the environmental noise affecting the QKD laser along its path. A key advantage is that this signal does not require strong attenuation (with a VOA) to the single-photon level; consequently, the resulting interference signal can be read by a classical photodiode that is characterized by a high Signal-to-Noise Ratio (SNR). This information is then used to actively correct the relative phase fluctuations between the twin fields [41].

The primary advantage of this scheme is its ability to satisfy both critical hypotheses required for a realistic TF-QKD implementation (Section 4.1.1). This capability stems from the combined use of two distinct lasers for comprehensive phase compensation:

1. The *reference laser* (ν_R), which is phase-locked to Alice's and Bob's local lasers (ν_A and ν_B) via a service fiber, guarantees the phase coherence of the twin fields at the source ($\nu_A = \nu_B = \nu_R$).
2. The *sensing laser* (ν_S), sent on the same fiber as the QKD signal via multiplexing, detects and compensates for phase fluctuations along the optical fiber.

Thanks to this dual stabilization mechanism, the Dual-Band Stabilization approach enables the achievement of a high Secret Key Rate using the TF-QKD protocol, even over optical distances of hundreds of kilometers where attenuation and phase fluctuations are substantially greater.

Experimental Results and Limitations

1) Optical Fiber Propagation Noise Cancellation. This new scheme demonstrates the ability to simultaneously perform key streaming via QKD and correct for phase noise originating from propagation in optical fiber (the second term of Equation (4.5)). This functionality allows for a more favorable duty cycle and a more efficient overall control of phase fluctuations.

Recent experiments under these conditions [41, 8] show how the solution facilitates the real-world implementation of the TF-QKD protocol, leading to significant increases in the coherence time of the interfering signals and reaching key generation rates that exceed 100 bit/s [41]. More generally, these experimental demonstrations have highlighted that the effectiveness of this scheme in correcting fiber noise is limited by a residual noise contribution given by [16]:

$$S_F(f, L) = \frac{(\lambda_S - \lambda_Q)^2}{\lambda_S^2} l L f^2 \quad (4.8)$$

Where λ_S and λ_Q denote the wavelengths of the sensing laser and the quantum signal, respectively.

NOTE According to this expression, we can indicate with $S_{F,A}(f)$ and $S_{F,B}(f)$ the terms that represent the resultant frequency noise for the two fiber paths, A and B. They are defined by substituting the respective path lengths, L_A and L_B , into the general function: $S_{F,A}(f) = S_F(f, L_A)$ and $S_{F,B}(f) = S_F(f, L_B)$.

The practical limit of this scheme's correction efficiency is directly related to the suppression factor $\frac{(\lambda_S - \lambda_Q)^2}{\lambda_S^2}$. This factor stems from the ideal assumption that the two lasers (sensing and quantum) interact identically as they propagate along the common optical fiber towards the central node, Charlie. In this perfect scenario, the sensing laser's phase information could fully correct the relative phase shift between Alice's and Bob's quantum fields without causing the collapse of the coherent states.

In a real-world scenario, however, the main limitation is that the two wavelengths (λ_S and λ_Q) are not perfectly equal but are only "sufficiently close" within the DWDM grid. As a result, they interact slightly differently with environmental noise, ensuring that a residual noise contribution always remains uncorrected. Furthermore, noise is introduced in the non-common fiber segments (such as those where the quantum laser modulation occurs). Since the sensing laser cannot detect this noise, it cannot be corrected, although this contribution is significantly lower than the total noise found in a simpler Common-Laser Scheme.

2) Residual Laser Noise Cancellation. A significant secondary benefit of this configuration stems from the phase-locking of the sensing laser to the reference source. Because the sensing laser is locked to the primary reference laser, the overall phase correction routine automatically suppresses the inherent phase jitter of the reference laser itself. More precisely, by exploiting this established phase relationship, the correction mechanism is able to treat the reference laser’s internal noise as a fixed offset, which is then eliminated during the phase comparison and compensation step.

Case Study: Dual-Band Stabilization in a Real Common Laser Setup

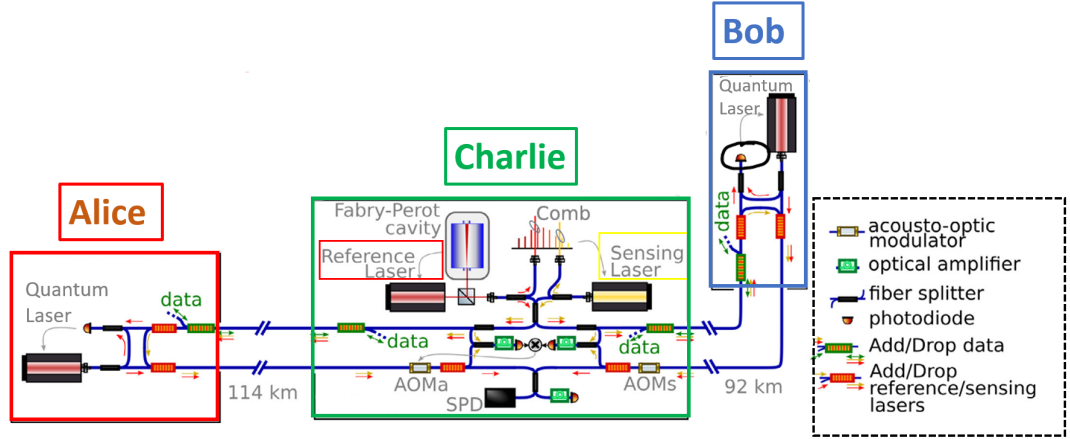


Figure 4.6: A real-world TF-QKD common laser scheme based on *Dual-Band stabilization*, implemented by INRiM [41]. This setup uses a reference laser (ν_R) and a sensing laser (ν_S) to stabilize the local QKD lasers. The sensing laser travels a round trip, first on a service fiber and then on a QKD fiber, accumulating phase information that is used in Charlie’s station to actively correct for fiber noise via an Acousto-Optic Modulator (AOM).

To better understand the operating principles and critical issues of the previously introduced theoretical solution, we will now examine a real-world implementation developed by the National Institute of Metrological Research (INRiM) in 2022, as shown in Fig.[4.6] [41]. The description of this scheme is particularly useful for the next chapters of this work, as it satisfies the critical hypotheses underlying a real-world TF-QKD protocol implementation.

First, it should be noted that the scheme uses two ultra-stable lasers for both the reference and sensing functions. These lasers have a linewidth of approximately

1 Hz and operate at standard frequencies of the DWDM grid. Specifically, the reference laser (ν_R) is frequency-stabilized using a high-finesse Fabry-Perot cavity, while the sensing laser (ν_S) is locked with a precise offset relative to the reference laser using an optical frequency comb. After being multiplexed in frequency, both lasers are sent to Alice and Bob through two separate service fibers.

At each node, the reference laser is extracted and used to phase-lock the local QKD laser. This QKD laser is then recombined with the sensing laser using DWDM, and both are finally sent back to Charlie on a separate fiber, which we will call the QKD fiber. This procedure is central to the proposed scheme because, since the sensing laser travels first with the reference laser on the service fibers and subsequently with the QKD laser on the QKD fibers, it contains all the necessary information about the accumulated phase variations of these propagation paths. This information can then be used for relative phase stabilization at Charlie's station.

Finally, the two distinct laser signals, now carrying their respective phase information, are simultaneously de-multiplexed and routed to specialized detectors based on their wavelength :

- The highly attenuated QKD laser is directed toward the single-photon detectors (D_0 and D_1). These detectors register the interferometric result of the two quantum lasers, which is subsequently processed publicly by Charlie and used by Alice and Bob to potentially extract a secret key bit.
- The brighter, unattenuated sensing laser is routed to a dedicated classical photodiode. This detector registers the interferometric result of the two sensing lasers and used this signal for continuous monitoring of the phase reference information required for the active phase alignment system.

The overall performance success of this scheme is fundamentally dependent on the efficiency of this phase stabilization system. An effective compensation scheme will significantly increase the duration of the quantum windows (τ_Q) and, consequently, increase the secret key rate.

However, we must account for the limitations of a real-world system. Because active stabilization cannot perfectly compensate for all external fluctuations, even with continuous phase referencing provided by the sensing laser, the phase compensation is never perfect over extended periods. Consequently, Charlie is still obligated to periodically interrupt the QKD transmission with *Classical window*. As in the original Common Laser Scheme introduced, he must introduce a dead time (τ_S) to perform a necessary, high-accuracy phase realignment using the photodiode. This periodic interruption inevitably introduces an operational duty cycle (d), which

quantifies the reduction in the final secure key rate.

A Real Scientific Experiment for TF-QKD with Dual-Band Stabilization

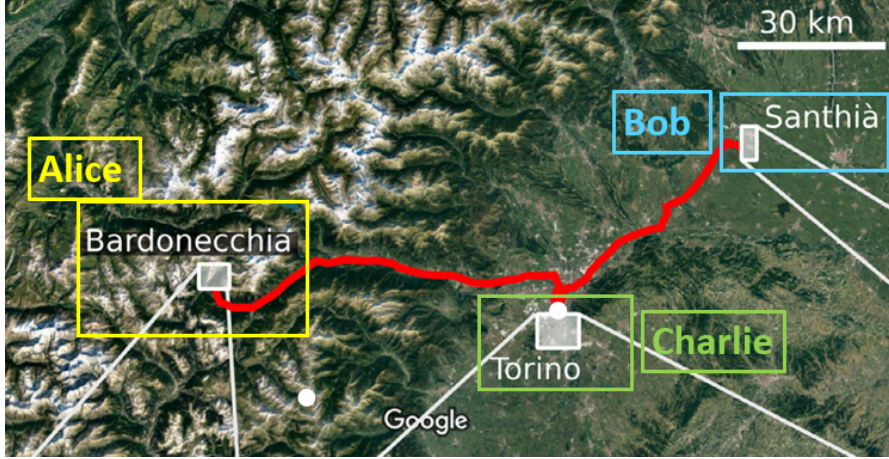


Figure 4.7: A schematic of the scientific experiment for TF-QKD using real optical fibers from the Italian Quantum Backbone [41]. Alice and Bob’s terminals are located in Bardonecchia and Santhià, and Charlie’s terminal is in Turin. This setup demonstrates the feasibility of dual-band stabilization for long-distance QKD over approximately 200 km, achieving a QBER of less than 1%.

The operating principle of the theoretical scheme shown in its most general form in Fig.[4.5] has been successfully applied in several scientific experiments [41, 45] to create practical architectures that solve the problem of relative phase fluctuations in the TF-QKD protocol. These implementations achieve this by distributing optical frequency between network nodes and applying active phase stabilization.

A notable example of the implementation of this stabilization is a scientific article published by the National Institute of Metrological Research (INRiM) [41], which used real optical fibers belonging to the Italian Quantum Backbone. Although the experiment did not perform a complete quantum key exchange, the results demonstrated the revolutionary potential of the Common-Laser scheme coupled with dual-band stabilization.

In particular, a compensation of the twin fields’ relative phase was achieved, maintaining the QBER below the 1% threshold over a 200 km distance. The crucial result of this study lies in the increase of the coherence time: in an unstabilized system (Fig.[4.3]), the QBER exceeds the critical 1% threshold in about $100 \mu\text{s}$. Thanks to dual-band stabilization (Fig.[4.5]), however, the quantum integration

time (τ_Q) was extended up to 100 ms (for 100 ms the phase is maintained with $\sigma_\phi = 0.13 \text{ rad}$, a QBER of only 0.5%). This extension of τ_Q by a factor of 1000 (from 100 μs to 100 ms) is fundamental for maximizing the secret key rate of the TF-QKD protocol.

A crucial step towards the large-scale adoption of TF-QKD.

We have now established that maintaining optical coherence between network nodes is the main challenge in implementing coherent quantum communications. The two solutions we've discussed have shown how this problem can be addressed in practical TF-QKD protocol implementation, increasing in this way, the effective duty cycle (d) of the QKD transmission.

Despite these advances, a significant challenge remains: the compromise between performance and the complexity of equipment and infrastructure. This is particularly evident in schemes like the one shown in Fig. [4.6], which, while offering excellent performance through ultra-stable lasers and high-finesse resonant cavities, have significant limitations in terms of cost, size, and robustness, making them unsuitable for commercial telecommunication networks.

A concrete solution to reduce implementation costs and optimize the use of TF-QKD equipment is to leverage a strong synergy between high-precision time/frequency distribution services and the creation of integrated network services for QKD. This approach is also supported by recent European initiatives such as the European Quantum Communication Infrastructure (EuroQCI) [46].

4.2.4 Independent Laser Scheme

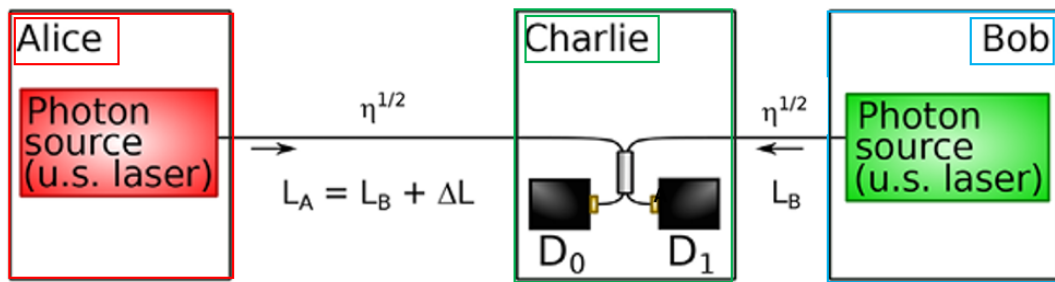


Figure 4.8: First proposed scheme based on the independent-lasers approach to TF-QKD, using ultrastable laser sources (u.s. lasers)[16]. This scheme shows the setup with Alice, Bob, and the central Charlie node, including their respective photon sources and detectors.

Introduction and Basic Analysis

This chapter introduces a new technique, the Independent Laser Scheme, which diverges from the previously discussed Common Laser Schemes. In particular, this method provides a foundational approach to stabilizing an optical channel's phase without a closed interferometer system, thereby achieving the required phase coherence for the TF-QKD protocol in a simpler manner. Furthermore, this technique is also applicable to multi-node networks with an optical channel asymmetry of hundreds of kilometers.

In this new scheme, the core idea is that two ultrastable lasers are phase-aligned once at the beginning of the quantum communication dedicated to key transmission, then allowed to evolve freely for a finite temporal interval before being realigned again. This topology is schematized in Fig.[4.8]

Similar to the derivation of Eq.(4.5) for the Common Laser scheme, the total phase noise [16] of the interference signal in Charlie can be expressed as:

$$S_{\phi}(f) = S_{l,A}(f) + S_{l,B}(f) + S_{F,A}(f) + S_{F,B}(f) \quad (4.9)$$

In this expression, the first two terms, $S_{l,A}(f)$ and $S_{l,B}(f)$, describe the intrinsic noise of the independent lasers used for TF-QKD at the two network terminals (Alice and Bob), while the last two terms, $S_{F,A}(f)$ and $S_{F,B}(f)$, represent the noise originating from the optical fibers that connect them to the central node (Charlie).

Independent Laser vs. Common Laser Scheme

Compared to the noise model for the Common Laser topology [Eq. (4.5)], it is noted that all terms related to the quantum signal propagation noise in the fibers of Eq.(4.9) ($S_{F,A}(f)$, $S_{F,B}(f)$) are reduced by a factor of four, appearing with a coefficient of 1 instead of 4 [16]. This significant difference stems from the fact that the previous scheme modeled a bidirectional optical path (a round trip). In real-world scenarios, and particularly when using parallel fibers within the same cable, the noise sources are highly correlated and accumulate coherently, justifying the original factor of four.

Unlike the previous model, this new topology's single-pass architecture leads to a different noise model. In particular in this scheme, the two lasers used by Alice and Bob are independent, and there is no round trip of radiation from Charlie to Alice (or Bob) and vice versa. Furthermore, it is also important to note that, although a service fiber may be used, its role is exclusively dedicated to classical communication during the authentication and post-processing phase and is not crucial for enabling quantum communication.

On the other hand, regarding the first two terms of Eq.(4.9) ($S_{l,A}(f)$, $S_{l,B}(f)$),

we can note that the coherence of the laser source has a direct impact on the overall performance of the phase correction system, in an analogous manner to the Common-Laser scheme. To extend the duty cycle and achieve greater phase stability, it is fundamental to use high-phase-coherence lasers. This can be achieved by using ultra-stable sources (such as those based on complex vacuum systems and very high-finesse Fabry-Perot resonators) or, alternatively, by adopting commercial narrow-linewidth lasers (Section 5.26).

This new proposed scheme offers several significant advantages, the most important of which are:

- **Simplified Network Topology and Infrastructure.** The scheme fundamentally simplifies the network by requiring only a single fiber for quantum key transmission between the communicating parties. This characteristic is critical for the scalability of complex, interconnected QKD networks.
- **Suppression of Back-Scattering Noise.** By eliminating the distribution of the reference laser through the service fiber, the noise component due to Rayleigh scattering is also removed. This prevents Rayleigh-scattered photons from propagating backward into the service fiber and then coupling into the QKD fiber through evanescent coupling. Such coupling would generate false clicks at the detector, thereby increasing the overall Quantum Bit Error Rate of the chosen QKD protocol.
- **Reduced Complexity and Cost of the setup.** Since a reference laser is no longer required, the need for extremely complex and expensive components for its stabilization—such as the high-finesse Fabry-Pérot cavity for laser stabilization (shown in Fig.[4.6])—is also eliminated. This represents a significant benefit for scalability, enabling the creation of more compact setups that are easier to integrate into modern telecommunication networks.

Nevertheless, the quality of the two laser sources in Alice and Bob continues to significantly affect the system’s final performance. This is because, as shown in Eq. (4.9), the relative noise terms of the lasers do not cancel out, even if the two channel lengths are perfectly balanced ($\Delta L = 0$), which did occur in the Common Laser scheme (see Eq. 4.5). Consequently, achieving greater overall stability and longer duty cycles in this new scheme requires either the use of lasers with high phase coherence or the implementation of strategies for compensating the laser phase misalignment in Charlie.

The primary disadvantage of this approach, therefore, is the stringent requirement

for not just one, but two, ultra-stable lasers, one for Alice and one for Bob. This means that if these ultra-stable lasers were still based on complex and expensive technologies, such as Fabry-Pérot cavities or vacuum systems, the overall setup remains challenging to integrate into a real-world telecommunication network.

4.2.5 Independent Laser Scheme with Dual Band Stabilization

Introduction

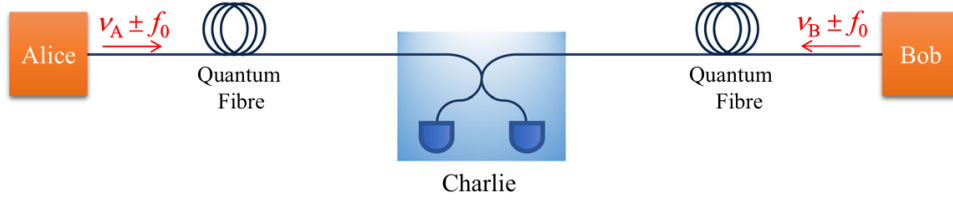


Figure 4.9: Scheme of the independent-lasers approach with dual-band stabilization for TF-QKD [16]. This method utilizes two ultrastable laser sources (ν_A, ν_B) and two sensing lasers ($f_{0A} = f_{0B} = f_0$) to enable high-bandwidth active phase noise cancellation on the optical channels from Alice and Bob to Charlie.

Despite the considerable advantages of the initial Independent Laser scheme, the most critical terms in Eq.(4.9) remain those related to the noise of the optical fiber. The main problem arises when extending TF-QKD to networks of a few hundred kilometers, where phase fluctuations is high and strongly dependent on environmental conditions. This makes periodic phase realignments necessary, which inevitably reduces the QKD duty cycle.

The best solution to overcome this limitation for long-range point-to-point connections is once again dual-band stabilization, as it can significantly reduce the need for phase realignment, allowing for very high duty cycles. To implement this, as in the case of the basic version of the Common Laser scheme, an additional laser (or, in this case, two different lasers) is added to the scheme. This auxiliary sensing laser travels on the same fiber as the single-photon signals sent by Alice and Bob, enabling high-bandwidth active phase noise cancellation for their respective optical channels (AC, BC). This improved scheme is shown in Fig.[4.9].

The key to this new scheme is to ensure that the two lasers—a sensing laser and a quantum laser—are coherent with each other at each node. By stabilizing the phase of the interferometer using the sensing lasers from the two nodes, the system also compensates for the quantum laser noise, thereby achieving the required phase

coherence for the TF-QKD protocol in a more straightforward manner. A technical demonstration of this procedure is reported in Section (5.3.2.)

Case Study: Dual-Band Stabilization in a Real Independent Laser Setup

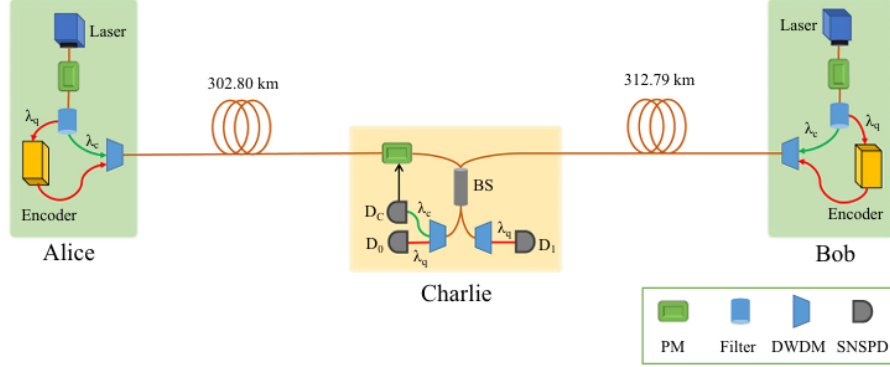


Figure 4.10: Real experimental setup used in [39] as an example of Independent Lasers with Dual-Band Stabilization. Alice and Bob, each equipped with an ultrastable laser (ν_A, ν_B), utilize a Phase Modulator (PM) to locally generate an optical frequency comb. Two spectral lines are extracted: the Quantum Signal (λ_q , encoded via the Encoder) and the Sensing Signal (λ_s , or sideband f_0). These signals travel through independent fibers to Charlie. At Charlie, the interference of the Sensing Signals ($\lambda_{sA,B}$) is detected by D_c to generate a real-time feedback signal. This signal drives a PID controller, which uses a PM on Alice's path to correct the phase drift of the reference channel, thereby automatically stabilizing the phase of the quantum channel ($\lambda_{qA,B}$) due to their fixed phase relationship. The corrected Quantum Signals then interfere, producing clicks on the Single-Photon Detectors (D_0, D_1).

The dual-band stabilization strategy, applied to the Independent Laser scheme, proves highly effective in minimizing the QBER caused by fast phase drift in the TF-QKD protocol. To demonstrate the feasibility and performance of this approach, we refer to the experimental setup presented in [39] and shown in Fig.[4.10]. This work served as a crucial case study, using a fiber configuration similar to MDI-QKD but without a service fiber. This confirmed the solution's ability to stabilize an open quantum link between distant, mutually incoherent nodes like Alice and Bob. The results reported in [39] concretely demonstrated the efficiency and robustness of the scheme, maintaining optical coherence over a distance of approximately 600

km. Crucially, with the Proportional-Integral-Derivative (PID) controller active, the interference output was precisely controlled, reducing the standard deviation of the relative phase to about 0.1 rad. This corresponds to a marginal contribution of only 1.6% to the total QBER of the TF-QKD protocol, confirming the validity of the independent laser scheme with dual-band stabilization even in scenarios with asymmetric fiber links.

4.2.6 Comparative Analysis and Critical Considerations

Common vs Independent Laser schemes with Dual-Band Stabilization

So far, we have presented two effective versions of two possible solutions—the Common Laser scheme and the Independent Laser scheme, both with dual-band stabilization—that can be adopted to solve the problem of phase drift of twin fields in a practical long-distance TF-QKD application. After having analyzed in detail the principles of both solutions and having presented real experiments based on them, such as those described in Fig.[4.6] and Fig.[4.10] respectively, it remains to compare the performance of the experimental setups proposed to date.

For this comparison, instead of simply considering the SKR of the open channel system (which is also influenced by other parameters such as detector performance or the chosen clock frequency for phase modulation), we will focus on the contribution to the total QBER (Section 2.2.3) that derives exclusively from the residual relative phase fluctuations between the twin fields that have not been correctly compensated.

Specifically, using the results reported in [39], which evaluates a series of real experiments that adopted the same specific version of the TF-QKD protocol, known as the SNS TF-QKD protocol, we can make a comparison based on the QBER contribution in the X basis and the Z basis (for this choice, see Section 3.5.3). From this simple analysis, we can deduce that the Independent Laser scheme with dual-band stabilization presented here provides a QBER of 4.75% (Fig.4.10) due to residual phase drift, a result only slightly better than other setups based on the Common Laser scheme with dual-band stabilization (Fig. 4.6), which present a QBER contribution of 5% due to uncompensated phase.

NOTE: It is important to emphasize that in this case, this is not a complete and formal analysis, since in a real QKD setup it would be necessary to also evaluate the impact of other contributions to the total QBER, as illustrated in Section (2.2.3).

However, for the purposes of this work, this result is more than sufficient, as it definitively demonstrates that a real setup based on Independent Laser schemes with dual-band stabilization not only leverages its advantages for the integration of

TF-QKD into today's global telecommunications infrastructures but also shows no degradation in QKD performance in terms of QBER caused by relative phase drift.

Main Criticalities and Possible Solutions

Having understood the potential of Independent Laser schemes with dual-band stabilization, it is now fundamental to analyze its main criticalities in order to address them.

The primary problem for a practical application of this solution, such as the one presented in Fig.[4.10][39], is due to the continuous employment of non-commercial components, such as ultrastable lasers. These lasers, often based on resonant cavities and vacuum systems, remain expensive and technically complex, making the technology unsuitable for the large-scale integration of TF-QKD protocols into existing network infrastructures.

For this reason, the next chapter will demonstrate an experimental architecture (Section 5.1) that uses this setup as a reference point for the techniques employed. More precisely, this new setup is improved to allow for the implementation of the independent laser scheme not with ultrastable lasers, but with conventional narrow-linewidth ones. This is made possible by leveraging the phase modulation mechanism for coherence transfer, in combination with dual-band stabilization.

Ultimately, it will be demonstrated that the independent laser scheme is the only one of the two approaches analyzed so far that allows for the creation of TF-QKD systems that are simultaneously compact, cost-effective, and easily transportable.

Chapter 5

Our Experiment

5.1 Introduction

This work focuses on the realization of a new experimental setup that implements the most effective topology described to date for stabilizing the relative phase between the twin fields in the TF-QKD protocol. For this purpose, the independent laser configuration with dual-band stabilization (Solution 2.1) was used, which represents the most promising approach for addressing the critical challenges outlined in Section (4.1.1) for the realization of a compact and transportable setup suitable for the real-world application of TF-QKD in modern telecommunication networks.

The present work can be considered an evolution of the 2022 experiment conducted at the INRIM laboratories [41], with a focus on implementing recently proposed phase stabilization techniques. By exploiting these advancements, we present a setup distinguished by its robustness and compactness.

5.2 Experimental Setup Overview

5.2.1 Architecture and Core Principles

This section describes the diagram of our experimental setup (Fig. [5.1]), detailing its architecture and fundamental operating principles. The physical implementation of this schema was realized at the INRIM laboratories in Turin, as presented in Fig. [5.2].

The implementation is based on a pair of distinct lasers at each terminal (Alice and Bob): a sensing laser ($\lambda_s = 1542$ nm) dedicated to phase compensation and a quantum laser ($\lambda_q = 1543.3$ nm) for the actual QKD key exchange. In each

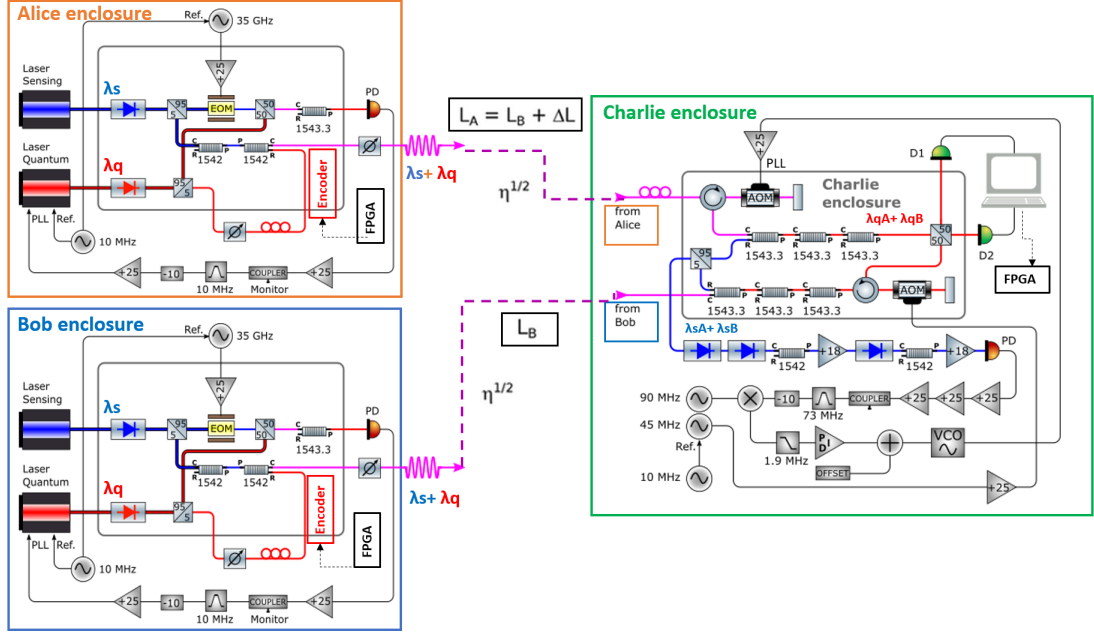


Figure 5.1: Diagram of our experimental setup, which addresses key challenges of long-distance TF-QKD in realistic fiber infrastructures. The architecture is based on an independent laser scheme with a dual-band stabilization strategy (Solution 2.1). It features a pair of lasers at each terminal (Alice and Bob): a sensing laser ($\lambda_s = 1542.1$ nm) for phase compensation and a quantum laser ($\lambda_q = 1543.3$ nm) for key exchange. These lasers are mutually phase-locked at each terminal using a PLL. The signals are then sent to a central Charlie node, which uses a third PLL to correct fast phase drift. The system's effectiveness is further enhanced by four active control loops for fast phase drift, slow phase drift, time realignment, and polarization mismatch compensation, ensuring a robust and secure protocol.

terminal, the QKD laser is phase-locked to the sensing laser, ensuring their mutual coherence. The quantum laser signals then pass through an encoder controlled by a dedicated FPGA board. These boards perform the decoy state and phase encoding required by the specific variant of the chosen TF-QKD protocol and the necessary attenuation with a VOA. The signals from both lasers are then combined using DWDM (Dense Wavelength Division Multiplexing) filters and sent to Charlie on two independent optical fibers (L_a and L_b).

At the central node, Charlie, the sensing and QKD laser radiation incoming from Alice and Bob are separated based on their wavelength. The sensing lasers are phase-locked to each other by a third PLL, which allows the detection and compensation of both the phase noise introduced by the fiber and the relative fluctuations of

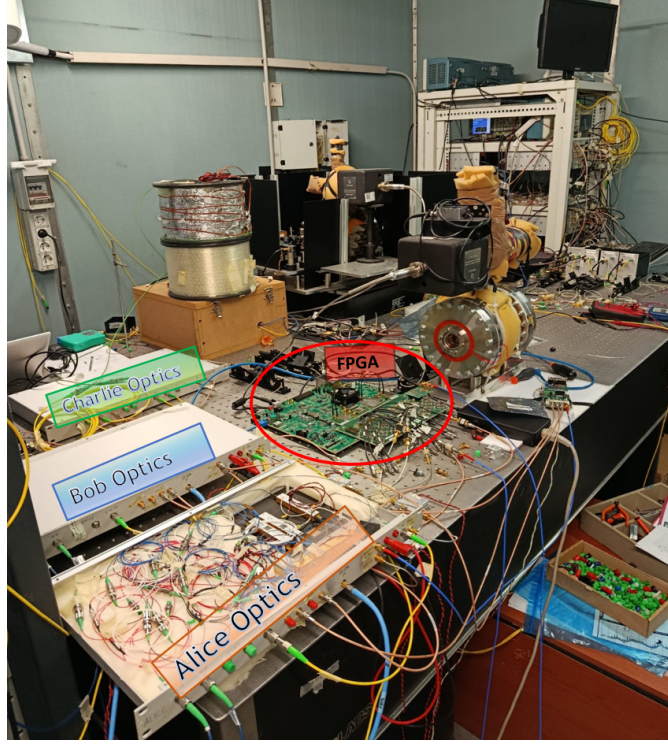


Figure 5.2: Portion of the experimental setup utilized for our TF-QKD experiment conducted at the INRIM laboratories in Turin. The three network nodes—Alice, Bob, and Charlie—were collocated in separate enclosures within a laboratory where the ambient temperature was actively stabilized at 23°C. The terminal nodes were connected to the central node via dedicated optical fibers.

the two sensing lasers in Alice and Bob. Since each sensing laser is already locked to its respective quantum laser, the correction applied to the sensing lasers also extends to the quantum signals, suppressing fiber noise and resolving any initial difference between the two sensing lasers.

5.2.2 Active Feedback and Control

To ensure the two quantum lasers interfere coherently at the Charlie node, the overall system coherence must be maintained. While a classical analysis requires monitoring the total optical power (P_{tot}) recorded during interference, for our quantum signals, the relevant quantity is the expected number of photons at the detector, $\langle n_{\text{out}} \rangle$, which is directly correlated to the probability of interference.

This is expressed as:

$$\langle n_{\text{out}} \rangle = \langle n_A \rangle + \langle n_B \rangle + 2\sqrt{\langle n_A \rangle \langle n_B \rangle} \cos(\Delta\phi) \cos(\Delta\theta) \quad (5.1)$$

In this context, $\langle n_A \rangle$ and $\langle n_B \rangle$ represent the expected number of photons per pulse of Alice's and Bob's signals recorded at Charlie (as described by the Poisson distribution, see Section 2.1). The term $\cos(\Delta\phi)$ is the interference contribution related to the phase difference ($\Delta\phi$) between the two quantum signals, and $\cos(\Delta\theta)$ accounts for the interference visibility loss due to polarization misalignment ($\Delta\theta$).

From this equation, we derive that the performance of our real-world TF-QKD experiment depends critically on multiple factors that are subject to environmental fluctuations. To mitigate these non-ideal factors and ensure that coherence is maintained, the setup integrates a system of four dedicated active control loops: two for phase compensation, one for temporal alignment, and one for polarization control.

Fast Phase Drift Correction

This high-frequency feedback loop operates on the sensing lasers to correct for rapid phase fluctuations (related to $\cos \Delta\phi$ term). The effect of phase fluctuations is the most dominant cause of visibility reduction, commonly referred to as "fast phase drift". This is corrected through an active feedback routine derived from those typically used in frequency metrology (such as Doppler cancelation). The core of this technique relies on three Phase-Locked Loops (PLLs) described in Section 5.3.

Crucially, the system implements this correction without the use of an additional service fiber, as it is based on the Independent Laser topology (described in Section 4.9). Furthermore, the fast phase drift can reach 30 rad/ms [41], which requires a feedback loop with a wide bandwidth.

Slow Phase Drift Correction

The fast drift correction system is inherently imperfect because the wavelengths λ_q and λ_s are different, causing them to interact slightly differently with fiber noise. Furthermore, non-common fiber segments (e.g., those dedicated to modulation at Alice's and Bob's terminals) introduce uncompensated phase accumulation (residual $\cos \Delta\phi$ term).

Consequently, a low-frequency digital control loop is implemented to correct this residual phase noise. This process requires the system to implement periodic re-alignment windows, during which the QKD transmission is interrupted. Although necessary, this introduces an inevitable dead time that reduces the protocol's SKR.

Overall, this mechanism works as a fourth, digital PLL that operates at a significantly lower frequency (between 50 and 100 Hz [39]) than that used for slow phase drift correction. It is described in Section 5.4.

Polarization Mismatch Correction

Dedicated strategies are implemented to mitigate polarization variations that would reduce the visibility of the quantum signal interference in the central node and thus increase the total QBER of the protocol.

For maximum interference visibility at Charlie, it is crucial that the quantum signals sent by Alice and Bob arrive with the same polarization. If the polarizations are not aligned, only their components along the same axes will interfere, causing a significant reduction in visibility, as indicated by the $\cos(\Delta\theta)$ term in the total power equation 5.25.

In the current setup, polarization control is performed manually using a polarization controller positioned at the end of the fiber connecting Bob to Charlie. This device is manually manipulated to achieve the condition of maximum interference, which corresponds to the optimal polarization alignment.

NOTE Ultimately, this process will be automated with the implementation of an active feedback routine based on an Electronic Polarization Controller (EPC), which can be driven by a voltage to continuously and precisely adjust the polarization. As seen in similar experimental setups [39], this routine typically operates at a continuous frequency of 5–10 Hz, ensuring the long-term stability of the system.

Temporal Delay Compensation

For two optical signals to interfere constructively or destructively, they must not only possess coherent phase and polarization, but they must also overlap perfectly in time. If one signal arrives late, their temporal overlap will be incomplete or absent, drastically reducing interference visibility and compromising quantum key quality. This makes the correction of any temporal delay crucial for maximizing interference at Charlie.

The primary cause of this delay is the difference in fiber length between the Alice-Charlie and Bob-Charlie paths. Although large length differences can be managed in a laboratory setting by adding dedicated fiber spools, this solution is not optimal for real-world deployments. In fact, this approach limits network flexibility and represents a significant obstacle for implementing TF-QKD in existing telecommunication networks.

The solution adopted in our setup addresses this issue by implementing a procedure managed by the FPGA boards at each node. The core of this process is to delay

the qubit encoding directly on the FPGA board that drives the phase and intensity modulators. This advanced system allows for the compensation of both minimal delays between the modulators within each node and significant delays due to fiber length differences between Alice and Bob. It is presented in Section 5.5

The ultimate goal is to ensure that the quantum pulses arrive at Charlie’s interferometer with a temporal superposition, within a very narrow 2ns window, to maximize the interference visibility.

5.3 Fast Phase Drift Cancellation System

This section provides a detailed description of the experimental setup of the three nodes utilized in our TF-QKD quantum communication system (Fig. [5.1]). In particular, it focuses on the implementation of a crucial active feedback system: the fast phase drift correction.

Our experiment utilizes a topology based on independent lasers with dual-band stabilization (Section 4.2.3). We will demonstrate the performance of this system in achieving relative phase stabilization between the twin fields in the TF-QKD protocol, even under challenging conditions such as:

- The two terminal nodes (Alice and Bob) are separated by hundreds of kilometers.
- The nodes do not a priori share mutual coherence.
- No additional service fibers are used, unlike schemes based on a Mach-Zehnder interferometer.

The primary goal of this efficient phase stabilization scheme is to suppress the most impactful noise terms in Eq. (4.9), which are primarily related to phase fluctuations in the optical fiber used for the exchange of quantum signals. These terms are particularly dominant in TF-QKD applications over networks hundreds of kilometers long, where the increased attenuation and phase fluctuations are highly sensitive to environmental changes (Section 4.1). Furthermore, we will demonstrate how the implemented scheme allows for the rejection of the intrinsic noise of the two quantum laser sources (Section 5.26).

5.3.1 Schematic of Alice and Bob’s Enclosure

In our experimental setup, each terminal node of the network (Alice and Bob) is equipped with two independent laser sources (Section 5.26), as illustrated in Fig. [5.3]:

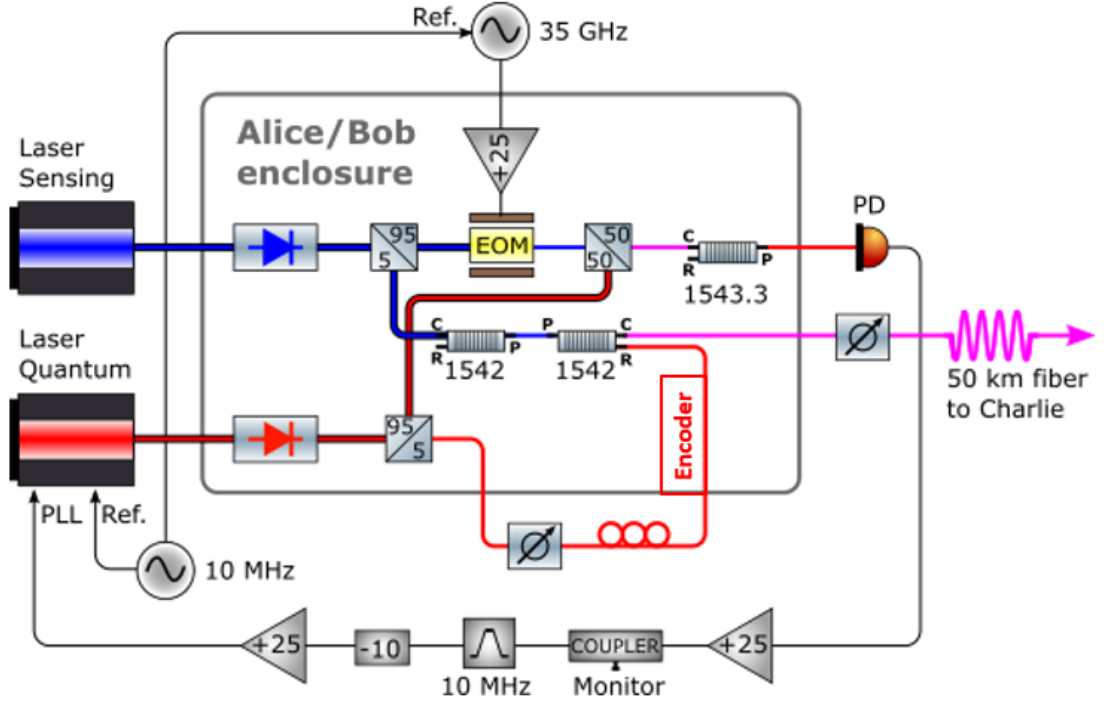


Figure 5.3: Schematic of the experimental setup for Alice’s (or Bob’s) enclosure, illustrating the local phase-locking mechanism.

- A **sensing laser** ($\lambda_s = 1542$ nm) used to detect fiber phase fluctuations (Section B.1).
- A **quantum laser** ($\lambda_q = 1543.3$ nm) which serves as the source for the quantum states (signal or decoy) used for key exchange and security (Section 3.5.3).

Our primary objective for the implementation of the fast phase noise cancellation system is to make these two spectrally separated lasers coherent through an active mechanism: the local Phase-Locked Loop (PLL).

The core challenge for this process lies in their large spectral separation, which makes it impossible to measure their beatnote directly. To overcome this obstacle, we exploit deep phase modulation of one of the two optical carriers (in our case, the sensing laser that follows a blue path in Fig[5.3]). This process generates high-order coherent sidebands that cover the spectral gap, allowing one of these sidebands to be beaten against the quantum laser to obtain a detectable frequency reference.

Principle of Phase Modulation

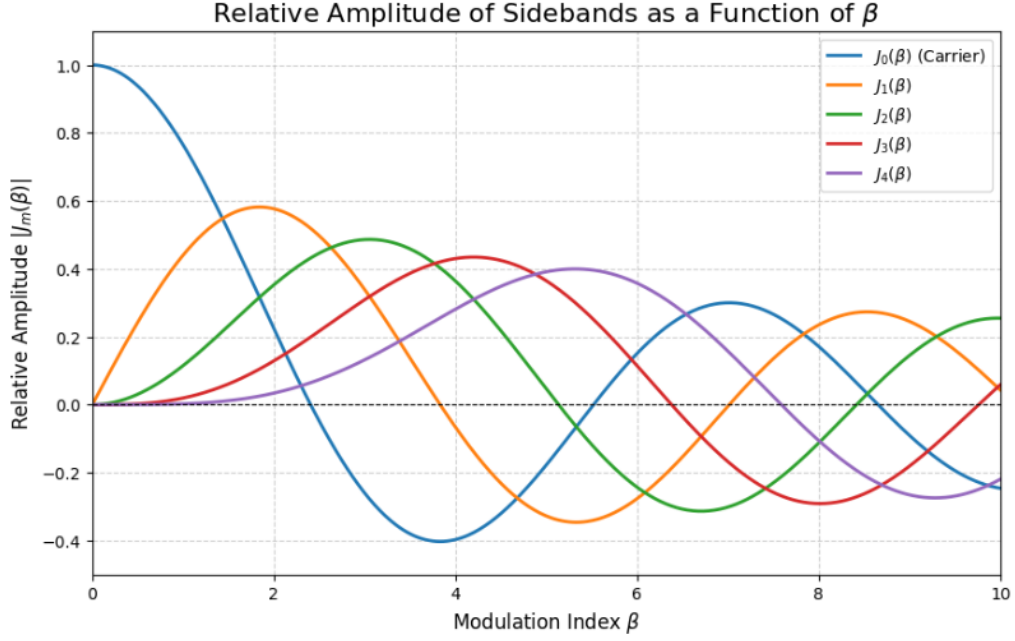


Figure 5.4: This graph illustrates how the energy of a phase-modulated signal is redistributed. As the modulation index (β) increases, the amplitude of the carrier frequency (J_0) decreases, while energy is transferred to higher-order harmonics ($J_1, J_2, \dots$). This physical principle means that with a higher modulation index, more discrete 'tines' of the frequency comb will appear and become more prominent.

To understand the underlying mechanism by which deep modulation of the sensing laser generates these high-order sidebands, a rigorous mathematical derivation will be presented below [47]. In particular, this demonstration will show how a modulated continuous-wave signal can be analytically rewritten as a summation of a series of discrete frequency components, which correspond precisely to the "teeth" of the generated frequency comb.

The modulation process begins with a continuous-wave optical signal, which can be represented generally as a monochromatic wave with a defined amplitude and phase:

$$E_{\text{in}}(t) = E_0 \cos(2\pi f_c t + \phi) \quad (5.2)$$

When this signal passes through a phase modulator (PM), driven by a radio frequency (RF) signal with frequency f_m , its phase is modulated. Consequently,

the resulting signal can be written in the form:

$$E_{\text{out}}(t) = E_0 \cos(2\pi f_c t + \beta \sin(2\pi f_m t)) \quad (5.3)$$

Where β is the modulation index, which depends directly on the amplitude of the applied RF signal, and f_m is the RF signal's frequency.

At this point, a series of purely mathematical steps is required to rewrite the previous expression.

Using Euler's formula, we can convert the cosine function into a complex exponential form, where:

$$\cos(x) = \frac{e^{jx} + e^{-jx}}{2} \quad (5.4)$$

Applying this formula to the output signal $E_{\text{out}}(t)$, we obtain:

$$E_{\text{out}}(t) = \frac{E_0}{2} \left[e^{j(2\pi f_c t + \beta \sin(2\pi f_m t))} + e^{-j(2\pi f_c t + \beta \sin(2\pi f_m t))} \right] \quad (5.5)$$

We can then separate the two terms:

$$E_{\text{out}}(t) = \frac{E_0}{2} e^{j2\pi f_c t} e^{j\beta \sin(2\pi f_m t)} + \frac{E_0}{2} e^{-j2\pi f_c t} e^{-j\beta \sin(2\pi f_m t)} \quad (5.6)$$

Next, we work on the term $e^{j\beta \sin(2\pi f_m t)}$ and expand it into a Fourier series using the Jacobi-Anger identity [47]:

$$e^{j\beta \sin(\theta)} = \sum_{n=-\infty}^{\infty} J_n(\beta) e^{jn\theta} \quad (5.7)$$

Where $J_n(\beta)$ are the Bessel functions of the first kind of order n . By substituting $\theta = 2\pi f_m t$, we obtain:

$$e^{j\beta \sin(2\pi f_m t)} = \sum_{n=-\infty}^{\infty} J_n(\beta) e^{jn(2\pi f_m t)} \quad (5.8)$$

Now we substitute this Jacobi-Anger expansion back into the signal equation (5.6) to obtain:

$$E_{\text{out}}(t) = \frac{E_0}{2} \sum_{n=-\infty}^{\infty} J_n(\beta) e^{j2\pi(f_c + n f_m)t} + \text{conjugate term} \quad (5.9)$$

The "conjugate term" is the complex conjugate of the first term. We can leverage the property of Bessel functions, $J_{-n}(\beta) = (-1)^n J_n(\beta)$, to combine the two complex terms and return to a cosine form:

$$E_{\text{out}}(t) = E_0 \sum_{n=-\infty}^{\infty} J_n(\beta) \cos(2\pi(f_c + n f_m)t) \quad (5.10)$$

This is the final equation of interest for our analysis.

This equation shows that the continuous optical signal, after modulation, is no longer composed of a single frequency (Fig.[5.4]). Instead, it breaks down into a series of discrete components. These components form the Frequency Comb, and they are separated from the carrier by an exact distance equal to the modulation frequency f_m and coherent to it.

Furthermore, the amplitude of each n -th "tooth" of the comb is determined by the corresponding Bessel function of the first kind, $J_n(\beta)$, where n is the harmonic order. In particular, the $n = 0$ term represents the carrier frequency, while the terms $n = \pm 1, \pm 2, \dots$ represent the sidebands. In Fig [5.4], the relative amplitudes of the various "teeth" of the frequency comb are shown as a function of the modulation index β . As can be seen, the amplitude of the carrier $J_0(\beta)$ decreases as β increases, eventually reaching zero, with the energy transferring to higher-order harmonics.

Generation of coherent frequency comb in our setup

The generation of the coherent frequency comb, which is the first step for the implementation of the PLL, begins with the sensing laser (λ_s). The majority of its optical power (95 %) is directed into an Electro-Optical Modulator (EOM), specifically a Phase Modulator (PM) (see the blue path in Fig.[5.3]). This component is driven by a radiofrequency (RF) signal at a frequency of $f_m = 35$ GHz, which is generated by a synthesizer locked to a stable 10 MHz reference to ensure high precision and stability.

The basic idea is to use the sensing laser as an optical source with a carrier frequency f_c corresponding to 1542 nm. When this signal passes through the PM driven at $f_m = 35$ GHz, according to the previous demonstration, it creates a frequency comb. As illustrated in Fig. [5.5], this comb extends over 140 GHz, with sidebands labeled from $m = 0$ to $m = 4$, equally spaced by the modulation frequency f_m .

A crucial detail of our implementation is the choice of which sideband to use for the PLL. Given the specific optical frequencies of the sensing laser (1542 nm) and the quantum laser (1543.3 nm), we observe that the resulting frequency separation is equal to 140 GHz. Therefore, to bridge this spectral gap and achieve phase locking, we determined that the optimal sideband for our system is the fourth harmonic ($m = 4$).

To ensure the desired $m = 4$ component of the frequency comb is sufficiently strong for detection, an appropriate modulation index β must be selected. Theoretically, maximizing the power of the $m = 4$ sideband requires a modulation index $\beta \approx 5.5$,

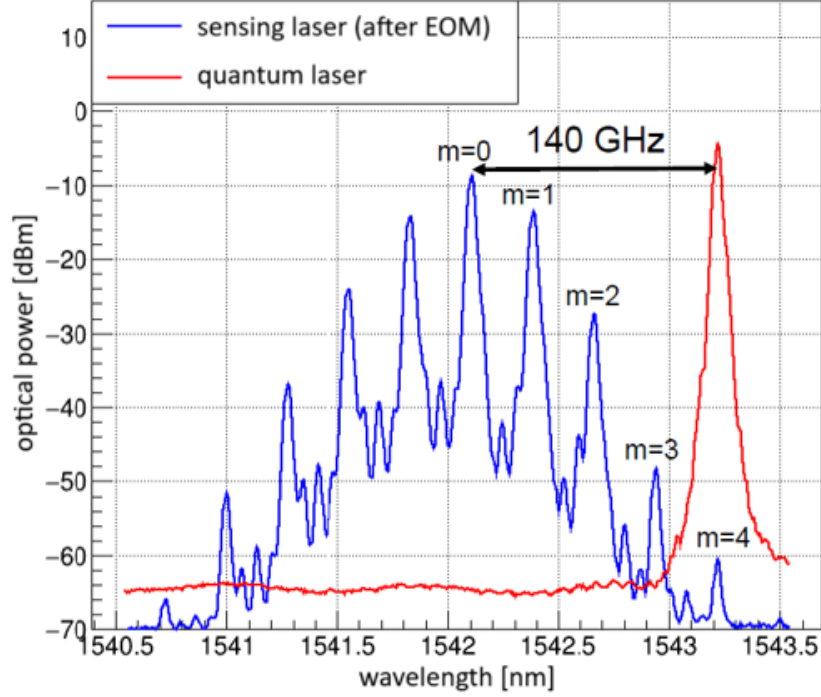


Figure 5.5: This figure shows the optical frequency comb generated by the NKT sensing laser after a phase modulator (PM). The comb extends over 140 GHz, with sidebands labeled from $m = 0$ to $m = 4$ that are equally spaced by the 35 GHz frequency of the modulating RF signal. The red line represents the quantum laser.

corresponding to the peak efficiency point for the fourth-order Bessel function (see Fig.5.4). However, practical limitations in the electronic driving circuit prevent us from reaching such high β values. Consequently, we utilize the strongest possible modulation signal available in our setup.

By observing the relative power of the sidebands, we estimated the actual operational modulation index. The power ratio between the $m = 4$ sideband and the carrier ($m = 0$) was found to be approximately 50 dB below the carrier. This power ratio is directly related to the amplitudes of the Bessel functions through the following equation:

$$\text{Ratio}(\beta)_{\text{dB}} = 20 \log_{10} \left(\frac{|J_4(\beta)|}{|J_0(\beta)|} \right) \quad (5.11)$$

By setting the observed ratio to -50 dB, we can estimate that the operational modulation index β in our setup is approximately 1.0. This value ensures a sufficient,

though not maximized, transfer of energy to the higher-order sidebands, making the $m = 4$ component suitable to serve as a locking reference.

Phase-Locked Loop (PLL) Implementation

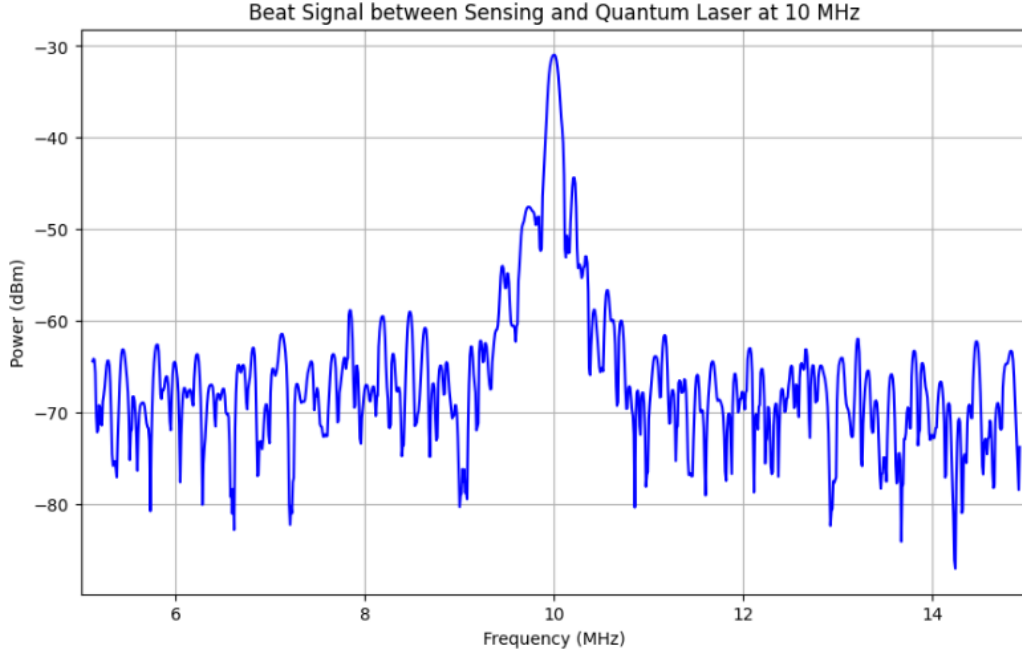


Figure 5.6: Spectrum of the beat signal between the quantum laser and the $m = 4$ sideband of the sensing laser, maintained stably at 10 MHz by the PLL.

To establish the necessary coherence, the quantum laser (λ_q) must be actively locked to the $m = 4$ sideband of the frequency comb generated by the sensing laser (λ_s). This process is executed via a PLL, a robust control loop designed to maintain a specific and stable frequency offset.

For this purpose, the optical power from the sensing laser and the quantum laser are interfered on a 50/50 beam splitter and directed to a photodiode (PD) with a bandwidth of approximately 10 GHz. The PD detects the beat note, which represents the frequency difference between the quantum laser and the $m = 4$ sideband (Fig. [5.6] shows the beatnote signal on a spectrum analyzer).

To close the control loop, the 10 MHz beat signal is amplified, band-pass filtered, and then phase-compared to a stable local oscillator at the same frequency using a mixer. A subsequent low-pass filter extracts the difference between these two signals, which yields the error term to be corrected. This error signal is then

processed by a Proportional-Integral-Derivative (PID) controller that generates an appropriate correction signal.

The described process is continuously iterated. At each cycle, the PLL uses the error signal to adjust the frequency of the quantum laser, thereby keeping the beat note constant at 10 MHz (Fig. [5.6]). The entire procedure can be managed via dedicated software (Fig. [B.2]), ensuring the beat frequency remains stable over extended periods.

Quantum Laser Encoder

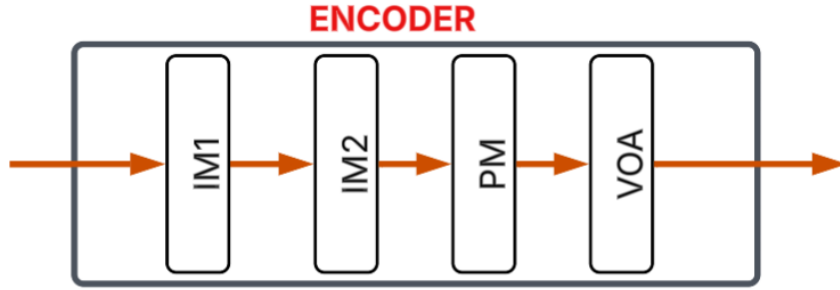


Figure 5.7: Encoder box: $IM_{1,2}$ = two intensity modulators in cascade; PM = phase modulator; VOA = variable optical attenuator.

During this operation, a small portion (5 %) of the quantum laser’s power is directed to the Encoder (see Fig. [5.3]).

The purpose of the encoder, schematized in Fig. [5.7], is to take a continuous-wave quantum laser signal as input and allow Alice (or Bob) to modulate it in both intensity and phase to implement the SNS TF-QKD protocol with decoy states. Specifically, the encoder receives the continuous signal and converts it into a pulse train with a frequency of 250 MHz and sets the intensity and phase of each pulse i according to the requirements of the chosen TF-QKD protocol.

Referring to Fig. [5.7], we can illustrate the main components of this encoder and explain its operation for the modulation of quantum signals.

The quantum signal arrives at the input via a polarization-maintaining (PM) fiber. Its polarization is aligned with the slow axis of a series of cascaded modulators: two intensity modulators (IM1 and IM2) and one phase modulator (PM). Each modulator is driven by a dedicated FPGA board; in our case, two identical boards were used for the Alice and Bob setups, respectively (Section 5.35).

- IM1 modulates the signal to create *signal* and *decoy* pulses.
- IM2 modulates the signal to create the same *signal* and *decoy* pulses as IM1, but its primary necessity is to ensure that the logical zeros correspond to a zero-photon level per pulse. This guarantees a sufficient extinction ratio (ER) for the cascaded modulator system when switching between the realignment and the quantum states, as a function of the chosen QKD communication length (Section 5.6.4).
- PM encodes the phase of each quantum signal pulse with one of 16 phase values ($\theta \in \{0, \pi/8, 2\pi/8, \dots, 15\pi/8\}$), satisfying the requirements for phase randomization and qubit encoding of the TF-QKD protocol.

Wavelength Multiplexing

After encoding, the quantum and sensing signals are combined using commercial optical Dense Wavelength Division Multiplexing filters (DWDM), typically featuring a channel spacing of 100 GHz. This technique allows the two signals, which are approximately 140 GHz apart (considering the sensing laser carrier and the quantum laser frequency), to achieve multiple critical functions simultaneously:

- They propagate on the same fiber.
- They are close enough in frequency to experience highly-correlated phase fluctuations along the fiber, enabling the use of the sensing laser fluctuations as a proxy for those experienced by the quantum signal.
- They are sufficiently separated in frequency to limit nonlinear effects, such as Raman scattering, which could generate inter-channel noise. This separation also ensures they can be effectively demultiplexed using conventional DWDM optical filters.

At this point, it is very important to stress that the sensing laser is phase-locked to the quantum laser via the frequency comb and the PLL described earlier. Consequently, since their absolute frequency difference is fixed and stable over time, any frequency variation of the sensing laser—caused, for example, by external noise resulting in a phase drift—will induce the exact same frequency shift in the quantum laser.

5.3.2 Charlie’s node Basic Schematic

The central node, Charlie, aims to measure and correct the total phase drift accumulated by the signals from Alice and Bob. This compensation is crucial as it

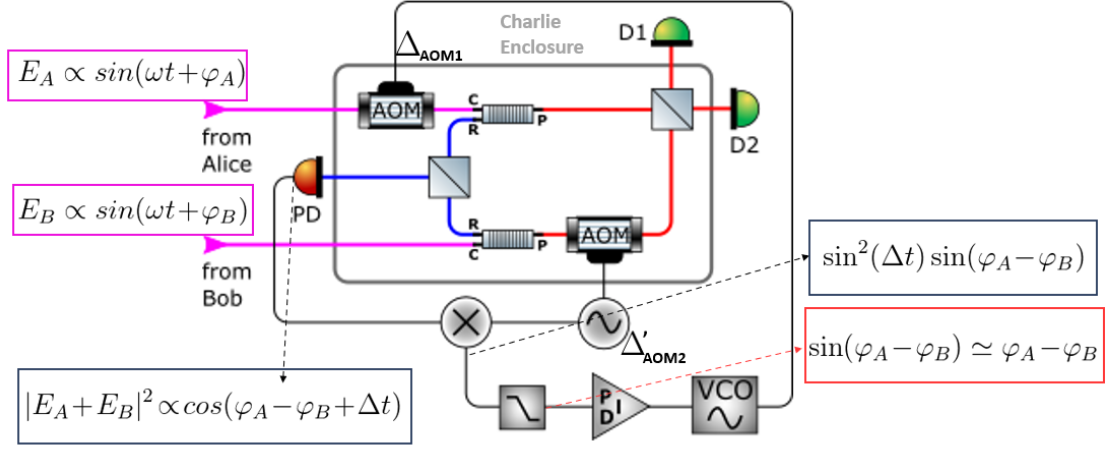


Figure 5.8: The diagram shows the basic schema of the central Charlie node's enclosure for a Twin-Field Quantum Key Distribution protocol. Its primary role is to measure and correct the total phase drift between the laser signals from Alice and Bob. The signals, with phases φ_A and φ_B , are combined, and a photodiode (PD) detects the beat signal. After being processed by a mixer and a low-pass filter, the output signal becomes proportional to $\sin(\varphi_A - \varphi_B)$. This output, which approximates the phase difference ($\varphi_A - \varphi_B$) for small fluctuations, is used to actively compensate for noise and maintain phase coherence, a critical requirement for the success of the QKD protocol.

corrects both the intrinsic noise of the two quantum sources and the noise induced by the respective optical fibers. This satisfies the two critical hypotheses of the TF-QKD protocol: that the twin fields are phase-coherent at the source and that they preserve this coherence along the path to Charlie (Section 4.1.1).

The following sections provide a conceptual and step-by-step analysis of this correction procedure.

Compensation of Phase Drift Due to Fiber Propagation

We assume that the sensing and quantum lasers initially share the same nominal frequency ($\nu_{s,A} = \nu_{s,B}$; $\nu_{q,A} = \nu_{q,B}$). However, upon arrival at Charlie's node, the lasers have undergone frequency shifts due to environmental noise ($\delta\nu$) experienced along their respective fibers (AC for Alice's link, BC for Bob's link). The frequencies arriving at Charlie are:

- Alice's signals: $\nu_{s,C_A} = \nu_{s,A} + \delta\nu_{F,AC}$ and $\nu_{q,C} = \nu_{q,A} + \delta\nu_{F,AC}$

- Bob's signals: $\nu_{s,C_B} = \nu_{s,B} + \delta\nu_{F,BC}$ and $\nu_{q,C} = \nu_{q,B} + \delta\nu_{F,BC}$

To enable the phase comparison, Alice's multiplexed signal is passed twice through an Acousto-Optic Modulator (AOM), denoted as AOM1, which is driven at a frequency Δ . This double pass introduces a frequency phase shift equal to 2Δ . Thus, the frequencies of Alice's signals incoming in Charlie are effectively shifted to:

$$\nu'_{s,C_A} = \nu_{s,A} + \delta\nu_{F,AC} + 2\Delta \quad (5.12)$$

$$\nu'_{q,C_A} = \nu_{q,A} + \delta\nu_{F,AC} + 2\Delta \quad (5.13)$$

Bob's multiplexed signal does not pass through any AOM in this stage, maintaining its original frequency shift due to the fiber noise.

The two sensing lasers ($\nu'_{s,C}$ and $\nu_{s,C}$) are then separated from the quantum lasers and combined on a beam splitter (BS) before being sent to a photodiode (PD). The PD measures the beat note frequency (ν_{beat}), which is the frequency difference between the two sensing lasers.

The measured beat note is equal to the sum of the AOM1-induced phase shift (2Δ) and the relative fiber-induced frequency error ($error = \delta\nu_{F,AC} - \delta\nu_{F,BC}$):

$$\nu_{beat} = |\nu'_{s,C_A} - \nu_{s,C_B}| = 2\Delta + error \quad (5.14)$$

To correct this total shift, a feedback loop is implemented using a radio-frequency mixer. The beat signal ($\nu_{beat} = 2\Delta + error$) enters the mixer, where it is combined with a reference signal from a local oscillator (LO) operating at the compensating frequency 2Δ . The mixer produces both the sum and the difference components of the two input frequencies. A subsequent low-pass filter extracts the difference term, which successfully isolates the true frequency error: $(2\Delta + error) - 2\Delta = error$.

In this manner, the relative frequency ($error$) between the two sensing lasers is successfully isolated for subsequent correction, despite the initial intentional frequency shift induced by AOM1.

Phase Drift Compensation for Sensing Lasers

The isolated error signal ($error$) is input to a PID controller, which generates a voltage signal proportional to the error term. The PID's output is connected to a Voltage-Controlled Oscillator (VCO), which translates this voltage signal into a correction frequency ($\delta\nu_{corr}$) applied to the AOM1 driver.

AOM1 applies this correction frequency, with the opposite sign, to Alice's signal. The total frequency shift applied by AOM1 is the sum of the nominal shift and the correction: $2\Delta + \delta\nu_{corr}$.

The objective of the PID is to drive the beat note frequency (ν_{beat}) back to the nominal reference value of 2Δ . This is achieved when the correction is equal and opposite to the noise-induced error, i.e., $\delta\nu_{corr} = -error$. This correction explicitly compensates for the differential noise between the two fiber links:

$$\delta\nu_{corr} = -(\delta\nu_{F,AC} - \delta\nu_{F,BC})$$

When the loop is closed and the system is correctly initialized, the final frequency difference between the two sensing lasers stabilizes at the fixed value of 2Δ without fluctuations:

$$\nu_{beat} = 2\Delta + error + \delta\nu_{corr} = 2\Delta + error + (-error) = 2\Delta$$

Once the loop is closed and the system is correctly initialized, the two sensing lasers from Alice and Bob are successfully phase-locked. This active feedback constantly corrects their absolute frequency difference, maintaining it at a fixed distance of 2Δ .

Consequently, Charlie's node successfully cancels the relative phase error between the two sensing lasers due exclusively to fiber-induced noise (under the initial assumption that $\nu_{sA} = \nu_{sB}$).

Role of AOM1 and Phase Drift Compensation for Quantum Lasers

The use of AOM1 is indispensable for this process of correcting the relative phase between the two sensing lasers. As the actuator of the PLL, AOM1 translates an electrical correction signal from the VCO into an effective optical frequency variation applied to Alice's signal.

However, AOM1 also affects the relative frequency between the two quantum lasers, which ultimately encode the key information. First, it is crucial to note that since AOM1 acts on the multiplexed laser output from Alice's node, it induces the phase shift 2Δ not only on the sensing laser, but also on the quantum laser. Second, the phase stability obtained from the PLLs in the terminal nodes ensures that each user's quantum laser is coherently locked to its respective sensing laser. Therefore, the phase correction applied to the sensing lasers in Charlie's node is reflected identically in the quantum lasers, since both undergo the same frequency shift due to fiber propagation.

The instantaneous frequencies of Alice's and Bob's quantum signals arriving at Charlie's node are:

$$\begin{aligned}\nu_{q,CA} &= \nu_{q,A} + \delta\nu_{F,AC_q} + 2\Delta + \delta\nu_{corr} \\ \nu_{q,CB} &= \nu_{q,B} + \delta\nu_{F,BC_q}\end{aligned}$$

where $\delta\nu_{F,AC_q}$ and $\delta\nu_{F,BC_q}$ represent the fiber-induced noise terms on the quantum link. The correction term, $\delta\nu_{corr}$, is derived from the sensing channel's error ($error = \delta\nu_{F,AC_s} - \delta\nu_{F,BC_s}$) and is designed to satisfy $\delta\nu_{corr} = -error$. Explicitly:

$$\delta\nu_{corr} = -(\delta\nu_{F,AC_s} - \delta\nu_{F,BC_s}) \quad (5.15)$$

The frequency difference between the two quantum lasers at the output of the control loop (after the correction, $\delta\nu_{corr}$) will thus be:

$$\nu_{q,C_A} - \nu_{q,C_B} = 2\Delta + (\delta\nu_{F,AC_q} - \delta\nu_{F,BC_q}) + \delta\nu_{corr} \quad (5.16)$$

Since the correction is the same as that applied for the sensing lasers and is designed to be equal and opposite to the sensing error, the final result is:

$$\nu_{q,C_A} - \nu_{q,C_B} = 2\Delta + \underbrace{[(\delta\nu_{F,AC_q} - \delta\nu_{F,BC_q}) - (\delta\nu_{F,AC_s} - \delta\nu_{F,BC_s})]}_{\text{Residual Error Term}} \quad (5.17)$$

The expression in brackets, which represents the residual error term, is non-zero because the fiber noise experienced by the quantum channel ($\delta\nu_{F_q}$) is not strictly identical to the noise experienced by the sensing channel ($\delta\nu_{F_s}$) due to the slight wavelength separation between the two bands. However, because both wavelengths propagate through the same fiber and the two laser bands are coherently locked within each node, this residual error term is minimized and considered negligible for practical stabilization.

The final frequency difference between the two quantum lasers stabilizes at a fixed value of 2Δ , free from the large fluctuations associated with the raw environmental error component.

$$\nu_{q,A} - \nu_{q,B} = (2\Delta + error) + (\delta\nu_{corr}) \quad (5.18)$$

Since the correction ($\delta\nu_{corr}$) is the same as that applied for the sensing lasers and is designed to be equal and opposite to the error ($\delta\nu_{corr} = -error$), the final frequency difference between the two quantum lasers stabilizes at a fixed value, free from the fluctuations associated with the error component. The final result is:

$$\nu_{q,A} - \nu_{q,B} = 2\Delta \quad (5.19)$$

Role of AOM2 and Final Correction

The final step is to eliminate this residual frequency difference, which is stable and equal to 2Δ . This is done by acting specifically on Bob's quantum signal. The multiplexed signal from Bob, which has not undergone any phase shifts from AOM1, is passed through a second AOM (denoted as AOM2). Its purpose is to induce a phase shift $2\Delta'$ on Bob's quantum laser. AOM2 is driven by a frequency

f_{LO2} , and the double pass induces a shift $2\Delta'$.

Now, the overall relative frequency difference between Alice's and Bob's quantum lasers becomes:

$$\nu_{q,A} - \nu_{q,B} = (2\Delta) + (-2\Delta') \quad (5.20)$$

When the shift induced by AOM2 ($2\Delta'$, equal to twice its driving frequency f_{LO2}) is made exactly equal to the reference shift of AOM1 (2Δ , equal to twice its driving frequency f_{LO1}), we get:

$$\nu_{q,A} - \nu_{q,B} = 0 \quad (5.21)$$

This is the final goal: to ensure $\nu_{q,A} = \nu_{q,B}$ before the lasers interfere and produce a result on the SPAD photodetectors.

The realization of this setup corrects the fast phase drift due to noise accumulated along the two fibers and satisfies the first critical hypothesis for a real implementation of the TF-QKD protocol: if the quantum sources share a mutual coherence, they will maintain this coherence even if they are hundreds of kilometers apart.

Mutual Coherence Between Terminal Nodes

This same correction process also allows us to satisfy the second hypothesis: that the terminal nodes of the network a priori share a mutual coherence. This is achieved by demonstrating that the total error corrected by Charlie's setup includes not only the relative phase fluctuations due to fiber noise but also the intrinsic frequency differences between the two quantum lasers.

This important property stems from the fact that Alice's and Bob's quantum lasers are phase-locked to their respective sensing lasers. Therefore, any initial frequency difference ($\Delta\nu_{intrinsic}$) between the quantum lasers also translates into the same frequency difference of the sensing lasers:

$$\nu_{q,A} = \nu_{q,B} + \Delta\nu_{intrinsic} \quad (5.22)$$

$$\nu_{s,A} \approx \nu_{s,B} + \Delta\nu_{intrinsic} \quad (5.23)$$

When the signals travel along the fiber, both undergo external noise. In Charlie's node, the beat note between the two sensing lasers reveals a total frequency error ($\Delta\nu_{total}$) given by the sum of the two noise components:

$$\Delta\nu_{total} = \Delta\nu_{fiber} + \Delta\nu_{intrinsic} \quad (5.24)$$

Since the intrinsic noise of the lasers and that of the fiber are considered statistically independent (uncorrelated), their standard deviations can be added in quadrature. Consequently, the PLL implemented in Charlie's node measures and corrects this single total error affecting the frequencies of the sensing lasers. Thanks to the

principle of phase locking guaranteed by the PLLs in the respective terminal nodes, the total correction performed by Charlie’s active control is applied identically to both quantum lasers.

The proposed solution let our system is able to simultaneously cancel both the fiber noise and the initial differences between the two quantum lasers, ensuring a stable interference for the QKD protocol. Finally, it is important to emphasize that this technique represents a notable improvement over phase correction schemes based on a Mach-Zehnder interferometer with a common laser source. Such schemes, in fact, required the use of a dedicated service fiber for the distribution of the reference laser’s frequency. The approach described here, instead, allows independent lasers to maintain coherence at a distance, by multiplexing the sensing and quantum signals on a single fiber, optimizing resources and simplifying the experimental setup.(Section 4.3).

5.3.3 Analysis of Phase Stability

To understand the results related to the stabilization of fast phase drift, it is essential to recall some fundamental quantities. In particular, the stability of the optical phase is a critical parameter in QKD systems, as it directly impacts the QBER (Section 4.1). Therefore, we quantify this stability by analyzing the phase variance (σ_ϕ^2), a key metric that measures the system’s performance and its corresponding QBER (e_ϕ) at the detector in the central node Charlie (Section 2.2.3).

Additionally, various noise contributions, each operating on a distinct timescale, influence the interference of quantum signals. To accurately identify these contributions, it is useful to analyze the system’s performance in the spectral domain. Specifically, we analyze the Power Spectral Density (PSD) of the phase noise, $S_\phi(f)$, in the frequency domain. In this way, it is possible to decompose the total phase noise into its different frequency components. Furthermore, we exploit the Wiener-Khintchine Theorem (Section 4.1), according to which the phase variance (σ_ϕ^2) is directly correlated with the integral of the PSD.

Experimental Configuration

In order to better characterize the performance of our fast phase stabilization circuit and evaluate its applicability in a real context, the laboratory configuration utilized optical fibers to separate the terminal nodes from the central node Charlie. Specifically, the links consisted of 15 m (Alice-Charlie) and 20 m (Bob-Charlie) of optical fiber. To emulate the losses of a long-distance link, an attenuator was included in each path, resulting in total link losses of 35 dB (Alice-Charlie) and 30 dB (Bob-Charlie), respectively. The total fiber length connecting Alice and Bob

was therefore 35 m, with a total attenuation of 65 dB.

More precisely, for this characterization procedure, we employed an alternative experimental setup instead of a true single-photon QKD configuration. This approach allowed for a more robust analysis of the system's performance. In particular, as shown in Fig. [5.9], an amplification stage was used to amplify the quantum laser signal, enabling the substitution of the Single Photon Avalanche Diode (SPAD) with a conventional photodiode (PD).

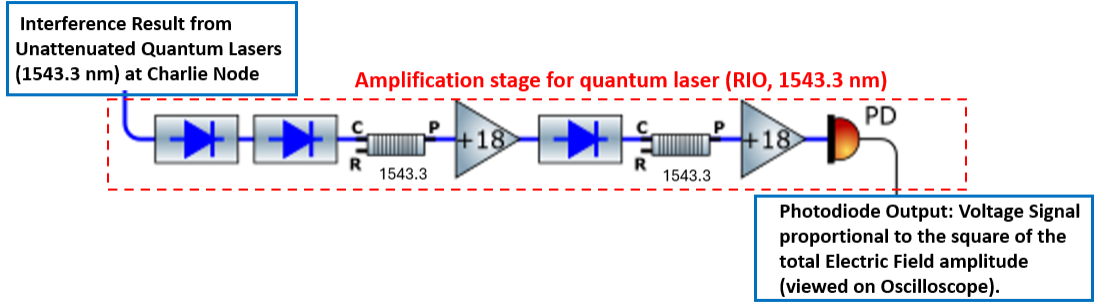


Figure 5.9: Diagram of the experimental setup utilized to evaluate the effectiveness of the phase stabilization circuit.

It is important to underline that with this new experimental setup, we are operating in a classical configuration. To ensure the two "quantum" (now classical) lasers interfere coherently at the Charlie node, the overall system coherence must be maintained. This requires monitoring the total optical power (P_{tot}) recorded at Charlie when the two signals interfere. Classically, this is expressed as:

$$P_{\text{tot}} = |P_A| + |P_B| + 2\sqrt{|P_A||P_B|} \cos \Delta\phi \cos \Delta\theta \quad (5.25)$$

In this context, $|P_A|$ and $|P_B|$ represent the optical powers of Alice's and Bob's signals recorded at Charlie. The term $\cos(\Delta\phi)$ is the interference contribution related to the phase difference ($\Delta\phi$) between the two signals, and $\cos(\Delta\theta)$ accounts for the interference visibility loss due to polarization misalignment ($\Delta\theta$).

NOTE In this classical configuration, the "quantum" signals interfering at Charlie's node generate a voltage proportional to the square of the total electric field amplitude. Consequently, this amplification and detection scheme provides a more stable and measurable signal for analyzing the phase stability circuit's performance.

Experimental Procedure

To definitively evaluate the effectiveness of our fast phase drift compensation scheme, we must detail the necessary steps for converting the voltage signal, as observed on an oscilloscope, into the cumulative phase standard deviation as a function of frequency. This process serves as the definitive metric for analyzing the system's performance.

1. Derivation of the Phase Equation

To begin our analysis, we first record the interference between the laser signals from the Alice and Bob nodes at Charlie's node. A photodiode converts this optical interference into a time-dependent voltage signal, $V(t)$.

The next crucial step is to establish a precise mathematical relationship between the measured signal of interference (e.g., the output voltage from the photodiode) and the phase difference ($\Delta\phi$) between the signals, as this enables the extraction of the instantaneous phase and the subsequent quantification of the phase noise.

Consider two laser signals, coming from the nodes Alice and Bob, whose electric field expressions (E_A and E_B) can be written as:

$$E_A = E_{0A} \cos(\omega t + \phi_A) \quad \text{and} \quad E_B = E_{0B} \cos(\omega t + \phi_B)$$

where E_{0A} and E_{0B} are the field amplitudes, ω is the angular frequency, and ϕ_A and ϕ_B are the initial phases. The total power measured by a photodiode is proportional to the square of the total field amplitude, $E = E_A + E_B$.

$$P_{\text{tot}} \propto |E_A + E_B|^2$$

By substituting the field expressions and neglecting the high-frequency terms that are not measurable by the photodiode, the average optical power (P) is:

$$P = P_A + P_B + 2\sqrt{P_A P_B} \cos(\Delta\phi)$$

Here, P_A and P_B are the powers of Alice's and Bob's lasers, and $\Delta\phi = \phi_A - \phi_B$ is the phase difference between the two signals.

A photodiode converts the optical power (P) into an electrical current (i). This conversion is governed by the photodiode's responsivity ($\mathfrak{R}$), such that $i = \mathfrak{R}P$ (where typically $\mathfrak{R} \approx 0.9 \text{ A/W}$). The current then passes through a resistance (R), generating a measurable voltage (V) according to Ohm's Law ($V = i \cdot R$).

The resulting signal voltage, representing the interference between Alice's and Bob's signals, is therefore given by:

$$V(t) = V_A + V_B + 2\sqrt{V_A V_B} \cos(\Delta\phi) \tag{5.26}$$

In this equation, V_A and V_B are the voltages associated with the individual signal powers (where $V_X = \Re RP_X$). An important observation to note is that this formula shows how the measured voltage is a combination of a constant term ($V_A + V_B$) and a term that varies based on the cosine of the phase difference ($\Delta\phi$).

To calculate the phase difference $\Delta\phi$, we need to rearrange the voltage equation. First, we can isolate the $\cos(\Delta\phi)$ term:

$$\cos(\Delta\phi) = \frac{V(t) - (V_A + V_B)}{2\sqrt{V_A V_B}}$$

Finally, by applying the arccosine function, we obtain the inverse formula that allows us to calculate the instantaneous phase difference $\Delta\phi$ from the measured voltage $V(t)$:

$$\Delta\phi(t) = \arccos\left(\frac{V(t) - (V_A + V_B)}{2\sqrt{V_A V_B}}\right) \quad (5.27)$$

This formula is often rearranged into a more practical form using the measured peak voltages. Let V_{max} and V_{min} be the voltages measured at fully constructive ($\Delta\phi = 2k\pi$) and fully destructive ($\Delta\phi = (2k + 1)\pi$) interference, respectively. In this way, the instantaneous phase difference can be calculated exclusively from the measured interference peaks as:

$$\Delta\phi(t) = \arccos\left(\frac{V(t) - \frac{V_{max} + V_{min}}{2}}{\frac{V_{max} - V_{min}}{2}}\right) \quad (5.28)$$

where $V_{max} = \frac{V_A + V_B}{2}$ and $V_{min} = \frac{V_A - V_B}{2}$.

2. Acquisition and Phase Derivation

The signal measured by the photodiode is the result of the interference between the unattenuated laser signals from the Alice and Bob nodes (Eq.5.25). This optical interference is converted into a time-dependent voltage signal, $V(t)$, which is proportional to the square of the total electric field amplitude.

At this point, it is possible to apply the derived Eq. (5.28) to the acquired data and converting in this way the raw voltage signal is into an instantaneous Phase vs. Time plot. The analysis of this plot is crucial for confirming the effectiveness of the active stabilization circuit.

In a non-stabilized condition, the relative phase would continuously drift due to various environmental noise sources. This drift is manifested as a problematic "folding" behavior (as observed, for example, in [41]), where the phase rapidly drifts across the limits of the normalized range (0 or 1), then "jumps" and restarts its

progression. Such large, uncontrolled oscillations would immediately indicate that the system is unusable for a QKD protocol.

Conversely, using the active stabilization circuit, the Phase vs. Time plot would exhibit the absence of a linear drift and significant large-scale fluctuations. The phase would be successfully maintained at a random oscillation around a fixed mean value (e.g., approximately 0.5), confirming that the coherence necessary for the TF-QKD protocol is successfully preserved [41]

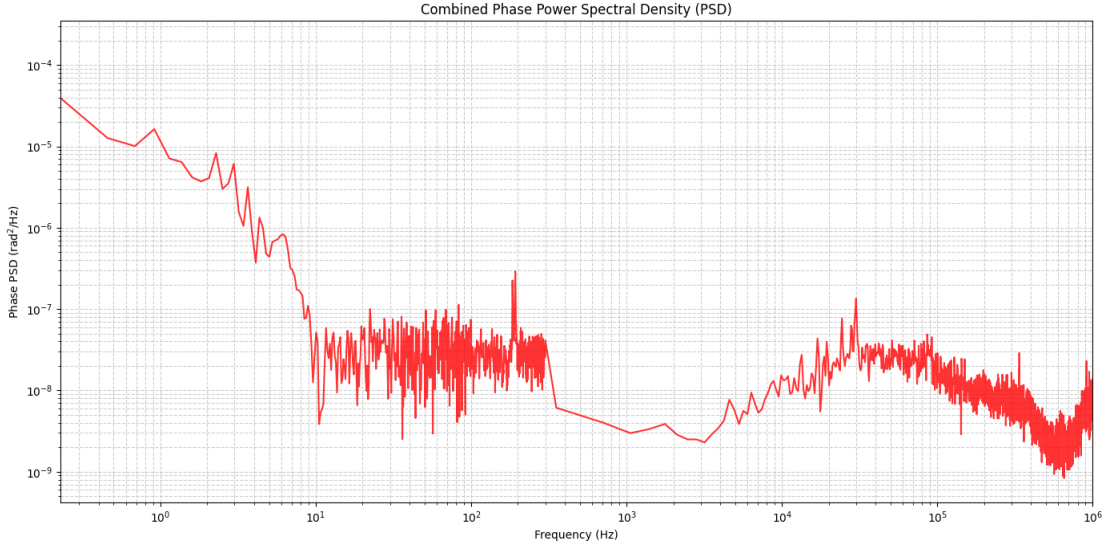


Figure 5.10: *Combined Phase Power Spectral Density (PSD)* of the stabilized phase. This graph was generated by combining two distinct PSD datasets, each calculated using the *Welch method* with specific parameters: a larger number of data segments for the high-frequency analysis ($f > 1\text{kHz}$) and a smaller number of segments for low-frequency analysis ($f < 1\text{kHz}$) to accurately capture the noise. The result is a single curve showing the stabilized phase noise distribution over a broad frequency range, from approximately 1 Hz to 1 MHz. The clear suppression of high-frequency noise confirms the stabilization circuit’s effectiveness in correcting fast phase drift. However, the residual low-frequency noise remains visible, which is likely attributable to the imperfect correlation between the noise experienced by the sensing and quantum signals due to the distinct fiber paths between them (see Section 5.4).

The analysis then moves to the frequency domain to understand the distribution of phase fluctuations. For this purpose, the PSD of the phase noise, $S_\phi(f)$, is calculated using the Welch method that allows to obtain a more reliable and uniform estimation of the noise distribution as a function of frequency.

The graph of the PSD of the stabilized phase noise, shown in Fig. [5.10]. Specifically, the analysis of this spectrum, shows that the phase noise is higher at low frequencies ($f < 10\text{Hz}$). This residual low-frequency noise is mainly attributable to the imperfect correlation between the noise experienced by the sensing and quantum signals due to the not common fiber paths between them (see Section 5.4). In contrast, for higher frequencies, the noise is suppressed by several orders of magnitude, which clearly demonstrates the effectiveness of the stabilization circuit in correcting fast phase drift.

NOTE It is crucial to emphasize that this measurement is performed by stabilizing the interferometer using the sensing laser while monitoring the quantum laser, leveraging the high correlation between the two bands established through the dual-band phase locking scheme.

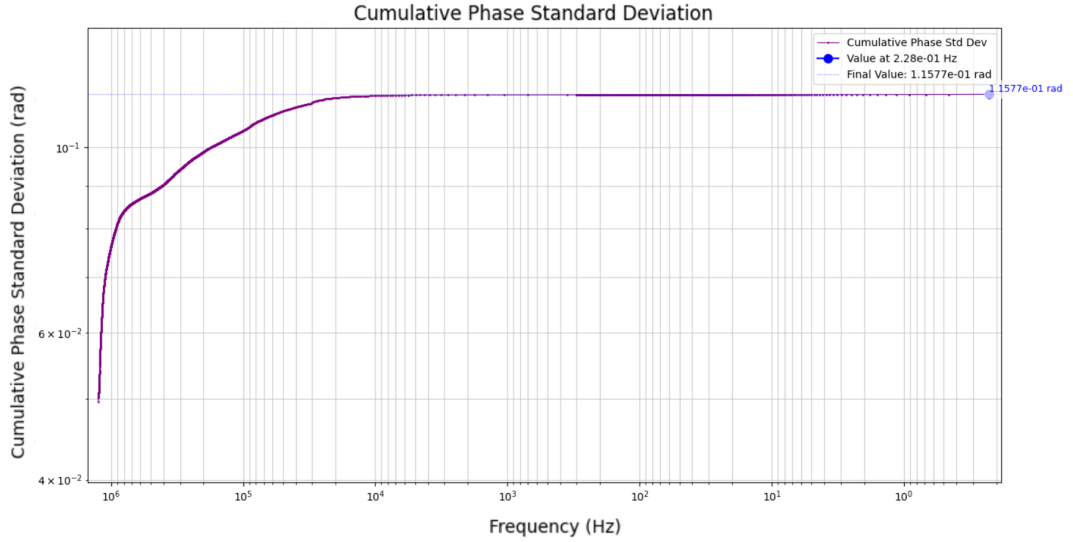


Figure 5.11: Cumulative Phase Standard Deviation (σ_ϕ). The graph shows the standard deviation of the phase between the two QKD lasers, calculated as a function of frequency. This result is obtained by numerically integrating the Phase Power Spectral Density (PSD), starting from high frequencies. The resulting curve represents the total contribution of phase noise across different timescales. The stabilization at approximately 0.12 rad demonstrates the circuit’s effectiveness in compensating for rapid phase fluctuations, ensuring the long-term stability required for the QKD protocol.

Finally, to evaluate the effectiveness of the stabilization circuit, all that remains is to quantify the total contribution of phase noise across the different frequency bands when the circuit is active.

For this purpose, the Cumulative Phase Variance (σ_ϕ^2) is calculated by numerically integrating the phase PSD ($S_\phi(f)$) from the highest measured frequency (f_{max}) down to a lower frequency defined by the inverse of the observation time, $1/T$. The cumulative variance, using the Wiener-Khintchine Theorem [48], is therefore given by:

$$\sigma_\phi^2(T) = \int_{1/T}^{f_{max}} S_\phi(f') df' \quad (5.29)$$

This integration operation allows us to determine the total phase noise variance contained in the frequency band ranging from $1/T$ to f_{max} . By iterating this process over a range of observation times (T), the result is the accumulation of the total phase noise in the frequency domain.

We thus obtain the final graph shown in Fig.[5.11], which represents the Cumulative Phase Standard Deviation (σ_ϕ) versus the Observation Time (T). This serves as the final metric for evaluating the stabilization circuit's effectiveness because it directly determines the maximum time interval (T) over which the system can maintain phase coherence below a required threshold (e.g., linked to a maximum acceptable QBER). Therefore, this analysis is essential for establishing the required periodicity of the realignment procedures necessary to compensate for residual low-frequency drift (see Section 5.4).

3. Analysis of Results

The graph in Fig.[5.11] shows the trend of the cumulative phase standard deviation (σ_ϕ) as a function of frequency, with the X-axis proceeding from high frequencies towards low ones.

At high frequency ($f > 10^4$ Hz), the value of σ_ϕ grows rapidly, indicating a significant contribution from high and medium-frequency noise. However, for very short observation times (less than 100 ms), the phase standard deviation remains below the critical value of 0.12 rad. This corresponds to a contribution to the TF-QKD protocol's QBER of less than 0.5% and confirms the circuit's effective action in suppressing fast phase drift.

Conversely, for frequencies below 10^4 Hz, the curve reaches and maintains a stable plateau at approximately 0.12 rad. This value represents the residual noise of our system for fast phase drift correction and corresponds to a QBER of about 0.5%. The stability of this value continues beyond 1 s of observation time.

These results demonstrate that the proposed dual-band scheme for phase stabilization is highly efficient. The setup guarantees a phase stability that is maintained for over four seconds, exhibiting a residual noise value of approximately 0.12 rad

and a corresponding contribution to the QBER of 0.5%. This result is fundamental, as it ensures that fast phase noise has a negligible impact on the total QBER for key exchange times exceeding four seconds.

Consequently, for the measurement interval considered, it will not be necessary to interrupt quantum communication for rapid phase realignments. This allows for longer duty cycles (d) and, as a result, a particularly high and reliable Secure Key Rate.

Our setup is therefore suitable for future integration of the TF-QKD protocol into modern long-distance telecommunication schemes, demonstrating excellent performance in phase stabilization while maintaining a portable, compact, and more accessible configuration compared to the solutions proposed in previously cited works (Section 4.3; 4.9).

Note on Results and Future Work It is important to underline that this result constitutes a simplified laboratory proof-of-concept. Since no tests were performed on spooled fiber or long-distance links, the short (35 m) fiber link utilized for this characterization serves as a preliminary test to benchmark and verify the ultimate limit of noise cancellation achievable by this dual-band scheme. These exceptional performance metrics, specifically the residual noise value of 0.12 rad, represent the best-case scenario. Therefore, the measurement will need to be repeated and validated in a real, long-haul optical fiber system to fully test its performance under practical conditions.

5.4 Slow Phase Drift Cancellation System

The rapid phase drift correction system is not capable of fully compensating for all fluctuations. This is due to several factors:

- The different wavelengths of the quantum and classical lasers (λ_q and λ_s) cause a slightly different interaction with the fiber noise.
- There are fiber segments not common to the two paths, such as those used for phase and intensity modulation in the Alice and Bob nodes.

These factors inevitably introduce a residual uncorrected noise, a phenomenon known as "slow phase drift." This type of drift typically requires periodically interrupting key transmission to perform a phase realignment, which creates a *dead time* and consequently reduces the Secure Key Rate of the protocol.

Based on the results from our simplified laboratory proof-of-concept, this slow drift is not visibly detectable on time scales shorter than 4 seconds. The graph of the

Cumulative Phase Standard Deviation (Fig. [5.11]) confirms this by showing a constant plateau at approximately 0.12 rad for an observation time up to 4.5 seconds. This indicates that, for the duration of our measurement, there was no further significant increase in noise due to slow drift.

NOTE This high stability is characteristic of our implemented short-link setup and contrasts with observations reported in other works (e.g., [41]) where, due to longer fiber links, significant slow phase drift is typically measured over similar time scales.

5.4.1 Experimental Setup and Analysis of Slow Phase Drift

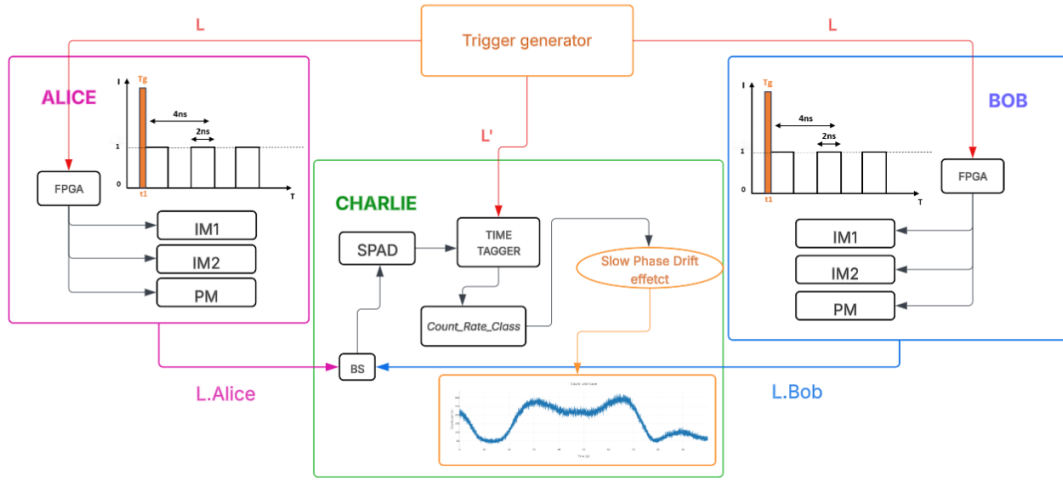


Figure 5.12: Experimental setup for evaluating the photon count rate over time. In this configuration, the interference of the quantum lasers from Alice and Bob is produced in the Charlie node, resulting in a click on the Single Photon Avalanche Diode (SPAD). The *CountRate Class* is used to process these clicks, producing a graph that allows for the precise evaluation of the effect of slow phase drift on the system.

To demonstrate the existence of slow phase drift and evaluate its impact on our system, we will work on longer time scales (on the order of hundreds of seconds) and modify the experimental configuration.

Instead of recording the interference between the QKD lasers on a photodiode (PD) as in Section 5.9, we will use a Single Photon Avalanche Diode (SPAD), located in the Charlie node. This modification allows us to transform the voltage signal as a function of time into a graph that now reports the photon count recorded on the

SPAD as a function of time.

The new setup, shown in detail in Fig.[5.12], uses a function generator to produce a square wave trigger. This wave is then sent to a Time Tagger and an FPGA to manage the synchronization process. The Time Tagger uses the trigger signal to open a time window (for example, 10 ms) within which the detector's clicks are counted. Simultaneously, the same trigger signal activates the FPGA which controls the phase and intensity modulation of the quantum signals generated by Alice and Bob and sent to the central node.

The advantage of this setup is that the interference of the two quantum signals at Charlie's node produces an event that is detected by the SPAD, but only the clicks recorded within the Time Tagger's counting window (with an appropriate delay, due to the difference in length of the fiber optic cables between the generator and the FPGA/Time Tagger) are considered as valid results of the quantum interference. Clicks outside of this window are discarded as noise.

The final output is therefore the detector's count rate, i.e., the number of valid clicks recorded as a function of time, taking into account the operating parameters of the SPAD. For example here we set the SPAD with a dead time of 10 ms and an efficiency of 10%.

Experimental Results

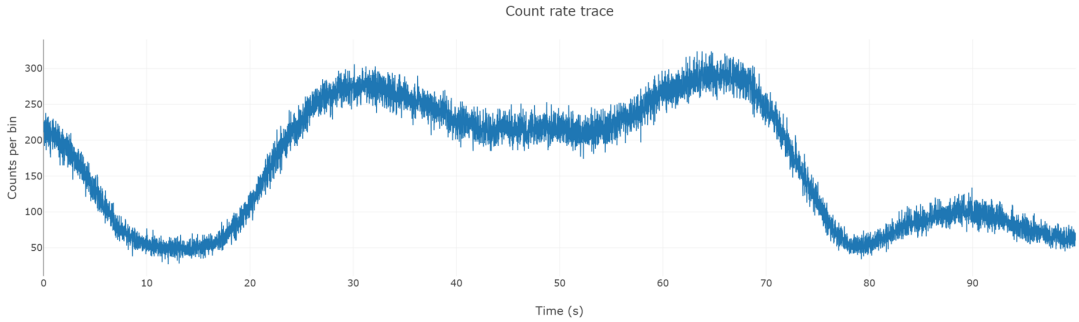


Figure 5.13: Photon count rate per 100 ms time bin, illustrating the effect of residual phase drift on quantum signal interference. In this scenario, the count rate is highly fluctuating (non-deterministic), deviating from the predictable maximum/minimum rates expected in an ideal deterministic condition. This uncontrolled slow phase drift makes the reliable implementation of the TF-QKD protocol impossible, underling the necessity of addressing this residual noise for secure key exchange.

By setting the temporal bin size to 100 ms, we obtain the graph shown in Fig.

[5.13]. This graph allows us to observe and quantify the effect of residual phase drift on the interference between the two quantum signals. To analyze this result, it's now important to distinguish two different conditions:

In a **deterministic condition**, i.e., with a well-aligned phase between Alice's and Bob's lasers, the result is predictable. For example, a phase difference of 0 or π rad leads to constructive or destructive interference, which translates respectively into a maximum or a minimum of the count rate for the detector considered in the 100 ms interval.

On the contrary, in a **non-deterministic condition** caused by slow phase drift, the count becomes random. The clicks occur stochastically with a 50% probability on each detector. Consequently, the count rate moves away from the ideal condition of maximum (or minimum) that would characterize a perfect alignment, proportionally to the extent of the drift.

The condition shown in Fig. [5.13] illustrates an intermediate scenario: the system's performance is neither perfectly deterministic nor entirely random. In particular, the phase stability is compromised by fluctuations, confirming that the regime is non-deterministic; however, the system is not entirely random because the fluctuations are bounded. In this scenario the average noise can still be quantified, although it is too large to maintain the required visibility for long time intervals.

It will therefore be essential to address the effect of residual phase drift in the system, since a non-deterministic interference condition makes it impossible to correctly implement the TF-QKD protocol for key exchange.

5.4.2 Strategy for Slow Phase Drift Correction

The correction of slow phase drift is based on the experimental setup already used in Section 5.3 for fast drift compensation. It specifically leverages the observation that the count rate of the single-photon detectors (D0, D1) is directly determined by the interference between Alice and Bob's quantum signals. This inherent relationship allows us to quantify the slow phase drift.

The procedure, therefore, essentially consists of monitoring the evolution of the number of counts over time. In particular, the ideal condition, with a well-aligned phase, corresponds to a maximum or minimum count. Any deviation from these ideal values, measured in a sufficiently short time interval so as not to lose alignment, generates a proportional correction signal. This signal is finally used for the application of the correction during the realignment phase, during which, we remember, the QKD transmission must be interrupted for a time that we want to be as small as possible.

The implementation of this process is achieved through a control loop, which acts as an additional PLL. This loop detects the deviation of the count rate from its maximum or minimum condition and generates a proportional error signal. This error signal is then input to a PID controller, which in turn generates the final correction signal. This correction signal is used to drive the AOM2 actuator (the one already used in the third PLL to correct the Δ shift of AOM1 on the quantum lasers. Section 5.3.2) in order to correct the slow phase drift during the realignment phase.

Overall, this technique implements a digital PLL for our setup for the periodic correction of slow phase drift. It is important to note that, unlike the feedback for fast drift, this system operates at a much lower frequency, between 50 Hz and 100 Hz.

Mathematical Model of the Re-alignment Strategy

Following the mathematical derivation described for fast phase drift compensation in Section 5.3.2, the difference between the instantaneous frequencies of the two quantum lasers, as observed at Charlie's node, is initially described by the following comprehensive equation:

$$\begin{aligned} \nu_{Q_{\text{Alice}}} - \nu_{Q_{\text{Bob}}} = & (\nu_{Q_A} - \nu_{Q_B}) + (\nu_{S_A} - \nu_{S_B}) + (\delta\nu_{\text{Fiber},Q,A} - \delta\nu_{\text{Fiber},Q,B}) \\ & - (\delta\nu_{\text{Fiber},S,A} - \delta\nu_{\text{Fiber},S,B}) + \delta\nu_{\text{AOM1}} - \delta\nu_{\text{AOM2}} \end{aligned} \quad (5.30)$$

Where the terms represent:

- $(\nu_{Q_A} - \nu_{Q_B})$: Intrinsic frequency difference between Alice's and Bob's quantum lasers.
- $(\nu_{S_A} - \nu_{S_B})$: Intrinsic frequency difference between Alice's and Bob's sensing lasers.
- $(\delta\nu_{\text{Fiber},Q,A} - \delta\nu_{\text{Fiber},Q,B})$: Difference in phase noise terms accumulated in the optical fibers for the quantum signals.
- $(\delta\nu_{\text{Fiber},S,A} - \delta\nu_{\text{Fiber},S,B})$: Difference in phase noise terms accumulated in the optical fibers for the sensing signals.
- $\delta\nu_{\text{AOM1}}$: Frequency shift introduced by AOM1, equal to twice its driving frequency ($2f_{\text{LO1}}$).
- $\delta\nu_{\text{AOM2}}$: Frequency shift introduced by AOM2, equal to twice its driving frequency ($2f_{\text{LO2}}$).

In an ideal dual-band stabilization system, several terms cancel out:

1. The first two terms ($\nu_{Q_A} - \nu_{Q_B}$) and ($\nu_{S_A} - \nu_{S_B}$) cancel due to the terminal nodes' PLLs, which phase-lock the quantum lasers to their respective sensing lasers (Section 5.3). This ensures that the two intrinsic frequency differences are equal.
2. The fiber-induced noise terms (third and fourth terms) are substantially mitigated by the high correlation between the sensing and quantum propagation noise, despite the wavelength separation.

Despite these compensations, the assumptions are not perfectly fulfilled, and consequently, a certain residual error still remains. This residual effect, is the slow phase drift that has not been compensated by the primary fast stabilization circuit and can be represented by an additional term, $\Delta\nu_{\text{Residual}}$.

Consequently, the frequency difference equation is simplified, representing only the terms that the slow drift circuit must manage:

$$\nu_{Q_{\text{Alice}}} - \nu_{Q_{\text{Bob}}} = (\delta\nu_{\text{AOM1}} - \delta\nu_{\text{AOM2}}) + \Delta\nu_{\text{Residual}} \quad (5.31)$$

The Proposed Solution To correct this residual drift, we must act on the frequency shift $\delta\nu_{\text{AOM2}}$ such that the net frequency difference between the two quantum lasers is driven to zero, thus reaching the phase alignment condition:

$$(\delta\nu_{\text{AOM1}} - \delta\nu_{\text{AOM2}}) + \Delta\nu_{\text{Residual}} = 0 \quad (5.32)$$

Isolating the necessary correction for AOM2, we obtain:

$$\delta\nu_{\text{AOM2}} = \delta\nu_{\text{AOM1}} + \Delta\nu_{\text{Residual}} \quad (5.33)$$

Recalling that the frequency shift introduced by each AOM is equal to twice the frequency of the respective local oscillator ($\delta\nu_{\text{AOM}} = 2f_{\text{LO}}$), we can substitute $\delta\nu_{\text{AOM1}} = 2f_{\text{LO1}}$ and $\delta\nu_{\text{AOM2}} = 2f_{\text{LO2}}$. The equation becomes:

$$2f_{\text{LO2}} = 2f_{\text{LO1}} + \Delta\nu_{\text{Residual}} \quad (5.34)$$

Dividing by 2, we obtain the final equation for the required correction frequency of the local oscillator for AOM2:

$$f_{\text{LO2}} = f_{\text{LO1}} + \frac{\Delta\nu_{\text{Residual}}}{2} \quad (5.35)$$

This mathematical demonstration indicates that the correction signal generated by the PID controller must be applied directly to the local oscillator frequency (f_{LO2}) of AOM2. The correction voltage, V_{corr} , generated by the PID controller's output,

must be precisely chosen to compensate for the residual error $\Delta\nu_{\text{Residual}}$, so that their sum is zero:

$$\Delta\nu_{\text{Residual}} + \Delta\nu_{\text{Correction}} = 0 \quad (5.36)$$

Consequently, when the system is in lock, the frequency of the local oscillator for AOM2 is appropriately modulated to cancel out the error, bringing the system back to the ideal alignment condition where $\nu_{Q_{\text{Alice}}} = \nu_{Q_{\text{Bob}}}$.

NOTE As shown in Fig. [5.8], in our setup, the local oscillator for AOM2 is a frequency synthesizer initially set to 45 MHz. The PLL must therefore be closed on this synthesizer, whose initial frequency of 45 MHz will be modulated for slow phase drift correction. In ideal conditions where $\Delta\nu_{\text{Residual}} = 0$, this 45 MHz frequency is exactly equal to the frequency produced by manually setting the adder that preceded the VCO with an appropriate voltage offset (Section 5.4). This ensures that the two frequency shifts ($f_{\text{LO1}}, f_{\text{LO2}}$) compensate each other perfectly, ensuring that the two quantum lasers achieve the same frequency.

5.4.3 Analysis of Slow Phase Drift Compensation

At this point, the fundamental questions for optimizing the slow phase compensation system are:

- How often do we have to perform the realignment?
- What is the duration of a single realignment cycle?

Regarding the first question, we established that our fast phase noise compensation circuit maintains a constant phase standard deviation for times greater than four seconds. This result strongly suggests that for time intervals shorter than four seconds, performing slow phase corrections is unnecessary, thus avoiding interruptions to QKD transmission.

Regarding the duration of the realignment, the goal is to find the minimum possible time. A shorter realignment duration is beneficial as it minimizes the reduction in the duty cycle for key exchange, consequently increasing the SKR of the protocol. However, the process must also be long enough to allow for an effective correction of f_{LO2} . This requires sufficient time to accurately measure the deviation of the photon count rate from its ideal value (maximum or minimum). In particular, since this process is based on gradient algorithms that exploit the slope of the count rate graph (Section 5.12) to evaluate the impact of phase drift, it is necessary to set an adequate time bin that produces a sufficient gradient for the algorithm to be reliably applied.

Given the inherent Poissonian statistics of photon counting, the uncertainty in the measurement (and thus the precision of the applied correction) directly depends on the number of photons detected. To minimize this uncertainty, the procedure requires either a high photon flux or a suitably long integration time. The photon flux is inherently upper bounded by the detector's limitations, such as the quenching effect, and cannot exceed a certain rate. As for the integration time, one must find the best compromise between measurement precision and minimizing the reduction of the system's duty cycle. The experimental results shown in Fig.[5.14] and

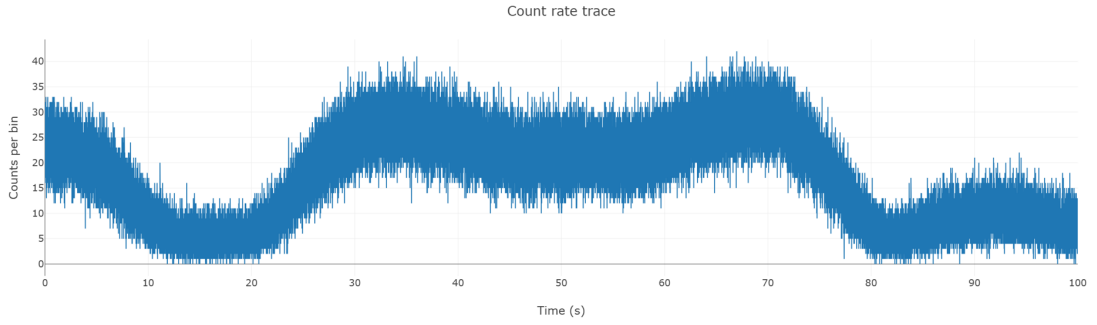


Figure 5.14: The evolution of the count rate for a temporal bin of 10 ms.

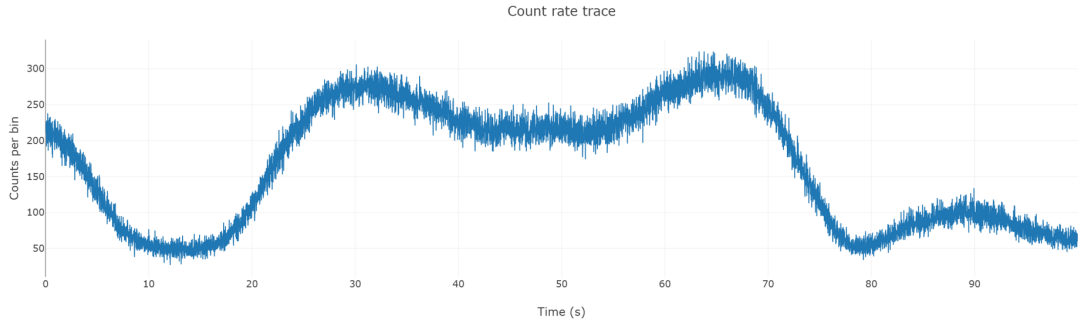


Figure 5.15: The evolution of the count rate for a temporal bin of 100 ms.

Fig[5.15] demonstrate the evolution of the count rate for time bins of 10 ms and 100 ms, respectively. These results highlight a crucial trade-off. While a longer temporal interval (e.g., 100 ms) extends the realignment time and reduces the duty cycle for key exchange, it also provides a significant advantage: an increased count rate and a better signal-to-noise ratio.

As a direct consequence of this improved ratio, the count rate curve for the 100 ms

interval (Fig. 5.15) is noticeably less noisy and shows a clearer trend compared to the 10 ms curve (Fig. 5.14). This less-noisy plot makes it significantly easier to successfully apply the gradient algorithm, which leads to a more effective and precise correction of f_{LO2} .

Conversely, reducing the acquisition time to shorten realignment results in a noisier graph, as seen in Fig. [5.14], due to a lower signal-to-noise ratio. For example, in the extreme case of an extremely small acquisition interval, such as 10 μ s, the signal-to-noise ratio becomes so low that the count rate exhibits a purely stochastic and unpredictable trend, making it impossible to apply a gradient-based correction algorithm.

In conclusion, an optimal strategy for slow phase drift realignment must balance these factors. Based on these results, performing a correction approximately every four seconds, with a temporal bin of around 100 ms for each realignment, is a good compromise to ensure a precise correction while minimizing key exchange dead time.

5.5 Temporal Realignment Procedure

The implementation of active feedback is essential for correcting the temporal delay between the quantum signals transmitted by Alice and Bob. This temporal realignment procedure is critical to ensuring the perfect superposition of the quantum pulses at Charlie's central node, a prerequisite for observing an interferometric measurement. The entire procedure is managed by dedicated FPGA boards at each node, compensating for path length differences through a sequence of delay adjustments.

The FPGA boards induce the appropriate delay by dividing the temporal correction into two sequential phases, relying on highly precise timing reference (for example, the Pulse Per Second (PPS) of the White Rabbit protocol [49] can be used). This guarantees the synchronization necessary for TF-QKD implementation.

Coarse Delay Calibration The Coarse Delay step compensates for large differences in fiber length between the Alice and Bob paths, which in a real network can be on the order of microseconds. As shown in Fig. [5.16], this delay is applied to synchronize the two paths. Our system can manage a coarse delay of up to approximately 250 μ s, which corresponds to a fiber length misalignment (ΔL) between Alice and Bob of about 50 km.

Fine Delay Calibration Once the coarse delay has been applied, the Fine Delay step performs a more precise adjustment to ensure the perfect temporal

```
[ ] # Add coarse delay to all modulators of Alice and Bob
    # The following code applies a coarse delay to all modulators of both Alice and Bob's nodes.
    # This ensures a preliminary time-alignment between the two nodes.

    # Coarse delay for Bob's modulators
    # delay_bob = 10x250ps
    base_gpio_ctr.write(df_b_IM1*4, 0 + delay_bob); # Apply coarse delay to Bob's IM1
    base_gpio_ctr.write(df_b_IM2*4, 5 + delay_bob); # Apply coarse delay to Bob's IM2
    # Phase
    base_gpio_ctr.write(df_b_PM*4, 7 + delay_bob); # Apply coarse delay to Bob's PM

    # Coarse delay for Alice's modulators
    # delay_alice = 19x250ps
    base_gpio_ctr.write(df_a_IM1*4, 0 + delay_alice); # Apply coarse delay to Alice's IM1
    base_gpio_ctr.write(df_a_IM2*4, 5 + delay_alice); # Apply coarse delay to Alice's IM2
    # Phase
    base_gpio_ctr.write(df_a_PM*4, 7 + delay_alice); # Apply coarse delay to Alice's PM
```

Figure 5.16: Code snippet used to add Coarse Delay to all modulators of Alice and Bob.

superposition of the pulses. This calibration is managed by writing to specific FPGA registers and also used for the internal alignment of the modulators within each terminal node (Alice and Bob), ensuring the pulses sent are coherent with each other before leaving the node. Regarding the fine delays, each delay step corresponds to 250 ps. The maximum manageable fine delay is 50×250 ps, which corresponds to a total delay of 12.5 ns. In terms of optical fiber length at 1550 nm, this is equivalent to:

$$L = v_{\text{fiber}} \times t = (204.2 \times 10^6 \text{ m s}^{-1}) \times (12.5 \times 10^{-9} \text{ s}) \approx 2.5525 \text{ m}$$

The combination of coarse and fine delay creates a robust and efficient calibration system, where the coarse delay handles large-scale compensations and the fine delay provides the necessary quantum precision refinements. In particular the overall realignment procedure consists of two sequential steps: internal modulator alignment and inter-node delay compensation.

1. **Internal Alignment (Node Level):** This procedure is performed by each terminal user (Alice and Bob) using their respective FPGA boards before initiating single-photon-level quantum communication. It compensates for the delays caused by the optical fiber within the three cascaded modulators in the encoder (Section 5.6.4). This step is crucial for ensuring that all pulses (signal or decoy) are correctly modulated before being sent toward the central node, Charlie. Specifically, each node's FPGA aligns its two intensity modulators to guarantee an Extinction Ratio (ER) of at least 40 dB (Section 5.6.4) and its phase modulator to ensure that the phase of each pulse falls within one of the discretized slices (Δ_k , Section 3.4.2) for the phase modulation.
2. **Inter-Node Alignment (System Level):** Once the modulators within Alice and Bob are internally aligned, the protocol proceeds to compensate for

the delay between the two nodes. This correction is carried out by the central node, Charlie, using its own FPGA board and Precision Time Protocol (for example the PPS of the White Rabbit protocol [49]).

The delay measurement proceeds as follows:

- All users (Alice, Bob, and Charlie) receive a synchronized PPS signal on their FPGA boards.
- Alice and Bob send a specific pulse train to Charlie (not dedicated to key exchange or phase realignment).
- Charlie, based on the PPS signal, derives the relative delay between the two users.

In principle, this procedure should be performed only once, in a preliminary manner, to allow each node to set the correct delays on its FPGA, based on the length misalignments initially observed in the fibers.

However, due to environmental noise, such as temperature variations, variations in the optical fibers can inevitably occur even during the implementation of the QKD protocol. Consequently, such realignment is required to be performed periodically, although with a much lower frequency than the other feedback routines.

5.5.1 Experimental Setup and Strategy

To verify the operation and effectiveness of our synchronization scheme, we used an experimental setup similar to the one for evaluating slow phase drift (Section 5.12). The original idea involves using a function generator to produce a square wave trigger, which simulates the PPS signal distributed by the White Rabbit protocol to the three network nodes. This wave is sent to a Time Tagger, located at the central node Charlie, and to the two identical FPGAs at the Alice and Bob nodes, to manage the synchronization process.

For our laboratory measurement, we implemented a simplified version of this scheme, using a single FPGA to manage both Alice and Bob and a single SPAD at Charlie's node. The basic idea remains the same: the trigger signal produced by the synthesizer is divided into two. One signal is sent to the single Alice and Bob FPGA, which uses it as a trigger to control the modulation of their quantum pulses. These pulses, generated, modulated, and sent by each node upon the arrival of the trigger signal, produce an event that is detected by the SPAD after their interference at Charlie's node. The other trigger signal is sent to the Time Tagger, serving as a START signal to open a time interval for counting valid SPAD clicks that fall within it.

using both coarse and fine delay corrections.

To implement the described temporal realignment procedure, the *Histogram* class

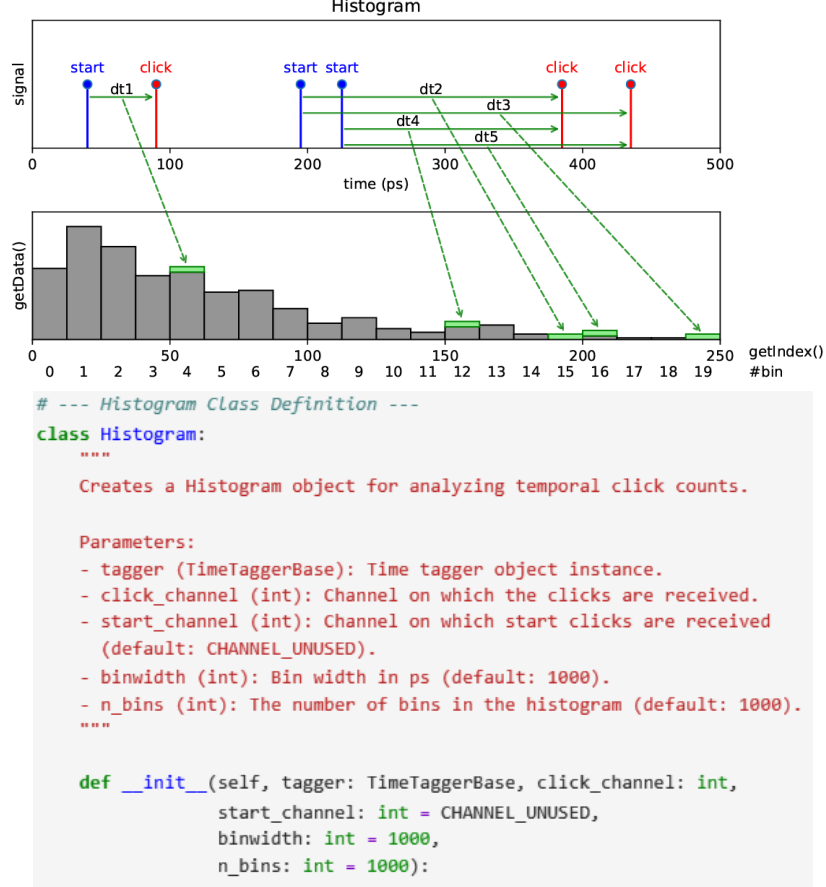


Figure 5.18: The Histogram class and its application for temporal synchronization in the experimental setup [50]. The down image illustrates how the class measures the time difference between a common *start signal* and multiple subsequent *click events*. In our experiment, the *start signal* simulates a White Rabbit’s PPS, while the clicks represent the *quantum pulses* detected by the SPADs. The up image shows the resulting histogram, which accumulates these time differences into bins. The two distinct peaks visible in the histogram correspond to the arrival times of the pulses from Alice and Bob, allowing us to precisely measure and correct the mutual delay between the two nodes using Charlie’s FPGA [50].

is particularly suitable for this purpose. Specifically, in our experiment, this class performs a measurement defined as *Multiple Start, Multiple Stop* (see Fig. [5.18 [50]). Its operates as follows: the system waits for a start event on a designated channel

(*start_channel*) and, for each start event, measures the temporal difference with respect to all subsequent click events that occur on another channel (*click_channel*). These temporal differences are accumulated in a histogram.

In the context of our experiment, the trigger signal generated and which simulates the PPS of the White Rabbit, serves as a common start signal for both Alice and Bob nodes; while the quantum pulses sent by each node, once detected by the SPAD, serve as click events. The resulting histogram will therefore show two distinct peaks that represent the "temporal position" of Alice's and Bob's pulses relative to the same start signal, allowing for the precise measurement of the mutual delay (see Fig.[5.18]).

The interval and resolution of the measurement are explicitly defined by the parameters of the Histogram class:

- **n_bins**: defines the number of time intervals of the histogram.
- **binwidth**: specifies the width of each bin in picoseconds, determining the temporal resolution of the measurement.

Events that occur outside the interval defined by the **n_bins** and **binwidth** parameters are ignored. This allows the analysis to focus only on the temporal window of interest, which in our case must be wide enough to include the two peaks of Alice's and Bob's events.

Steps for Temporal Realignment

The procedure for temporal realignment in our experimental configuration is divided into a series of conceptual and technical steps designed to ensure the perfect temporal overlap of quantum pulses at Charlie's node. In particular, we can identify three main phases:

1. Optimization of Intensity Modulator (IM) DC Points The initial objective is to optimize the operating position of each intensity modulator. The DC points of the Alice and Bob modulators are gradually adjusted while the phase modulators are kept off. The effectiveness of this operation can be verified in real-time by observing an increase in the peaks on the Time Tagger histogram as the DC points approach their optimal position, which corresponds to a "minimum logical zero" (close to zero photons) and a "maximum logical one" (total signal transmission).

This procedure is fundamental to ensuring that the intensity modulators operate at their point of maximum efficiency. An example of a successful optimization is shown by comparing the un-optimized peak in Fig. [5.19] with the optimized one in Fig. [5.20].

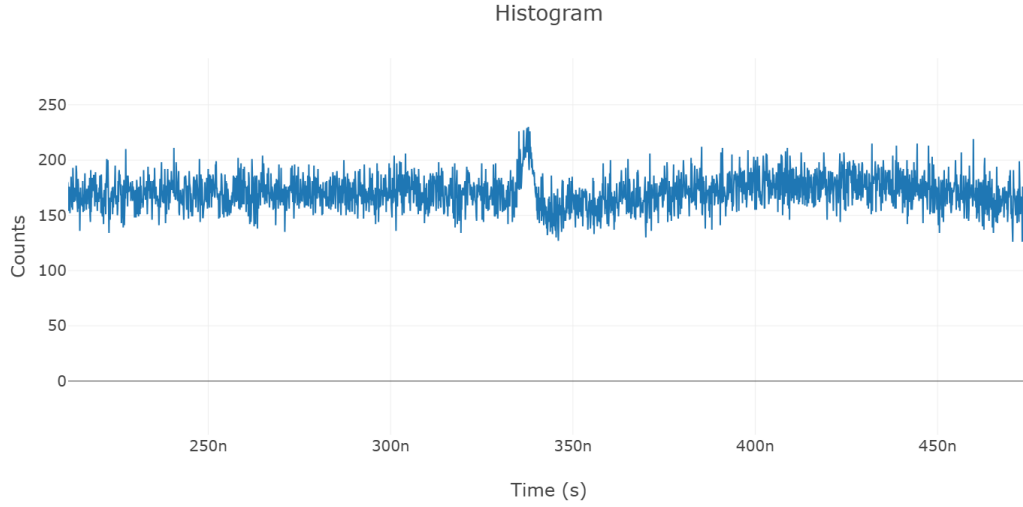


Figure 5.19: Histogram showing detected pulse counts before DC point optimization. The broader, lower peak (around 340 ns) corresponds to the optical pulses sent by Alice. This indicates a non-optimal DC bias setting for the intensity modulator, resulting in reduced signal transmission and a less clearly defined pulse at Charlie’s detector.

2. Compensation of the Delay Between Nodes Following the optimization of the intensity modulators, the temporal alignment between the Alice and Bob nodes is performed. Using the same START signal that simulates the PPS, two distinct peaks are observed on the Time Tagger histogram. The temporal difference between these two peaks represents the delay that needs to be compensated. This is accomplished by operating on Charlie’s FPGA, where coarse and then fine delays are applied to ensure perfect temporal superposition of the pulses.

3. Alignment of the Phase Modulator (PM) Once the Alice and Bob pulses are temporally aligned, they generate an interference pattern. The phase modulator (PM) of each node are then turned on to maximize the visibility of this interference pattern.

It is important to note that the temporal realignment procedure for our setup is based on a simplified but effective strategy:

- The alignment of the intensity modulators (IM1,2) and the phase modulator (PM) within each node is performed once, in a preliminary phase, using an oscilloscope. We assume this alignment remains constant because the length

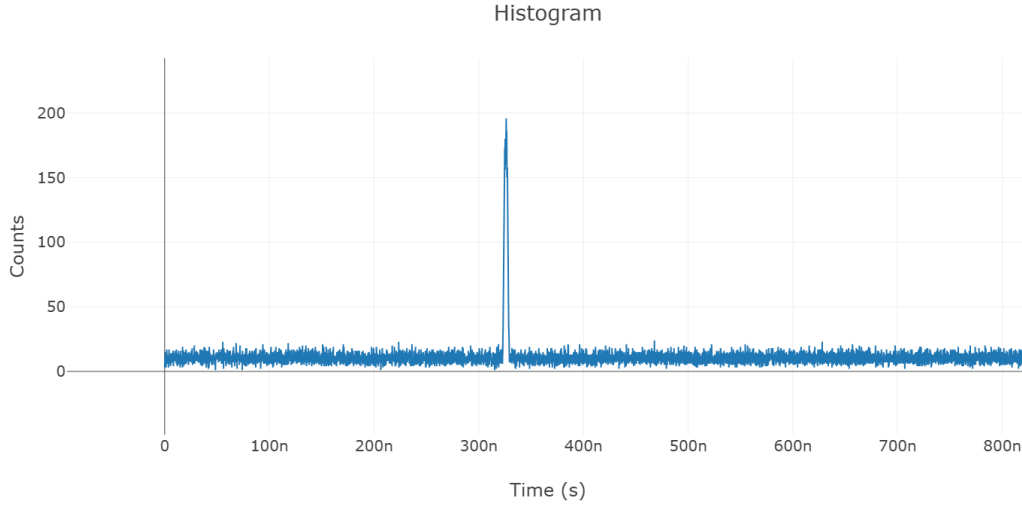


Figure 5.20: Histogram showing detected pulse counts after DC point optimization. The sharp, high peak (around 340 ns) demonstrates the successful optimization of the intensity modulator’s DC bias at the Alice node. The temporal delay remains the same due to the fixed fiber length between Alice and Charlie, but the improved DC point results in a clearly defined, high-transmission pulse. This is crucial for subsequent temporal alignment and confirms that the system is operating at its point of maximum efficiency.

of the internal fibers is subject to minimal variations.

- In contrast, the search for the optimal DC points for each modulator and the periodic compensation of the delay between the nodes (Alice and Bob) are performed before each key exchange session using the SPAD and the Time Tagger histogram at Charlie’s node.

Real-World Example of the Realignment Procedure

Let us now implement a real example of time delay alignment based on the previously described setup and code script.

For this specific measurement, the total acquisition time was set to 1 s, with the histogram updating dynamically with the data collected every 100 ms. This update interval allows for observing the evolution of the peaks during the search for the optimal DC points. The width of each bin was set to 100 ps and kept constant for the entire duration of the acquisition. The number of bins (`n_bins`) is consequently

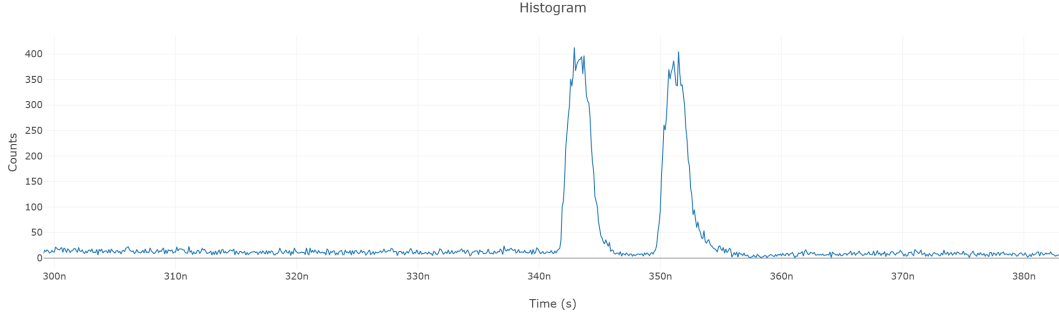


Figure 5.21: An example of a histogram for temporal realignment, showing two distinct peaks that represent the arrival times of pulses from the Alice and Bob nodes. The histogram was acquired over approximately 1 s, with a bin width of 100 ps. The separation between the peaks indicates the temporal delay that needs to be compensated using the FPGA in Charlie, which applies both *coarse* and *fine delay* adjustments. In this specific case, the measured delay between Alice and Bob is approximately 8 ns, which, considering the speed of light in a 1550 nm single-mode (SM) fiber (2×10^8 m/s), corresponds to a difference in fiber lengths (L_{A-C} and L_{B-C}) of approximately 1.6 meters. The width of each peak is approximately 2 ns, consistent with the expected duration of the modulated quantum pulses and, furthermore, it also confirms the correct temporal alignment of the cascaded intensity modulators (IM1,2) at their respective nodes, which are set with a fine delay. In this process, the characteristics of the detector, a SPAD, must be considered: in our case, the detector has an efficiency of 10% and a dead time of 10 ms after each click; however, potential Raman noise was not optimized in this particular measurement.

calculated as the ratio between the update time and the width of each bin:

$$n_{bins} = \frac{\text{Update Time}}{\text{Bin Width}} \quad (5.37)$$

By inserting the specific values from our measurement, we obtain:

$$n_{bins} = \frac{100 \text{ ms}}{100 \text{ ps}} = \frac{100 \times 10^{-3} \text{ s}}{100 \times 10^{-12} \text{ s}} = 10^9$$

These parameters must be defined at the start of our main function. Subsequently, the process involves acquiring time tags from designated channels in regular intervals, binning the arrival times, and finally, computing the histogram.

NOTE In this process, it is also fundamental to consider the characteristics of the detector used, a SPAD. In our case, it has an efficiency of 10% and is characterized

by a dead time of $10\text{ }\mu\text{s}$ after each click, making it insensitive to further photons in this interval. To avoid data loss, the external synthesizer was set to produce a square wave at a frequency of 50 kHz , ensuring a distance of $20\text{ }\mu\text{s}$ between two consecutive triggers, a value greater than the detector's dead time.

The result of this measurement is reported in Fig.[5.21]. The peaks observed on the histogram are expected to reflect the modulation parameters discussed in Section 5.6.4 for the quantum signals. Since in our setup the pulses have a duration of 2 ns , the peaks of Alice and Bob should have a width corresponding to this value.

5.5.2 Automation of Realignment with Cross-Correlation

Once the histogram is obtained, showing two distinct peaks for Alice's and Bob's pulses, the alignment procedure can be automated. This because while the peak height increases with optimal alignment, manually identifying the correct delay can be subjective. To make the measurement more objective and automated, an algorithm based on autocorrelation was implemented.

Autocorrelation is a powerful mathematical tool that analyzes a signal's intrinsic periodicity by measuring the similarity between the signal and a time-delayed copy of itself. When a signal has peaks that repeat at regular intervals, the autocorrelation produces a peak corresponding to that interval. In our specific case, the peak corresponds to the relative time delay between the two quantum nodes.

NOTE Autocorrelation is a special case of cross-correlation, as shown in Figure [5.22], where the functions F and G are identical.

In this specific application, the algorithm focuses on a defined region of interest within the histogram, which contains the event peaks for Alice and Bob.

After calculating the autocorrelation, the code identifies the first significant peak, intentionally ignoring the central peak at zero lag. In fact, the central peak represents the signal's perfect correlation with itself, which is not useful for measuring a relative delay between two distinct events. On the contrary, the position of the first significant peak away from the center provides an objective, numerical measure of the precise delay between the two nodes.

Analysis of Autocorrelation Results

The effectiveness of the automated temporal realignment procedure is best illustrated by comparing the autocorrelation results under two different experimental

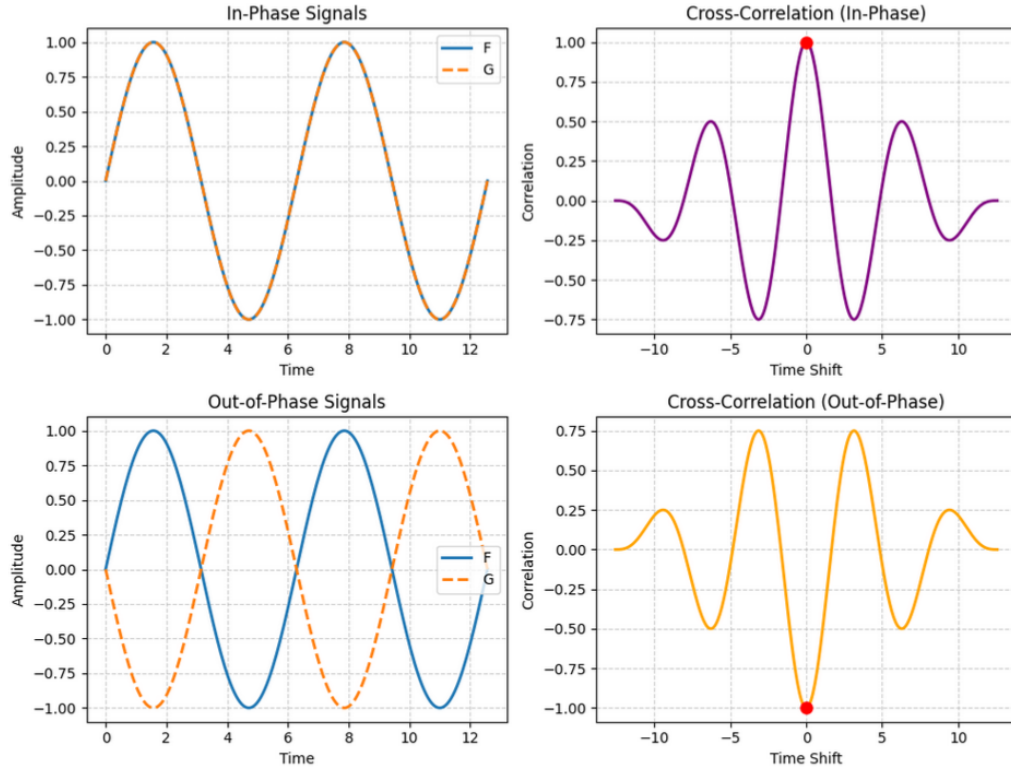


Figure 5.22: Illustration of *Cross-Correlation* for signal analysis. The four subplots demonstrate the correlation between two time signals, F and G. The top-left plot shows situation in which the signals are in-phase, while the top-right plot shows the corresponding cross-correlation result, with a red dot indicating the peak correlation. Similarly, the bottom-left plot shows the situation in which signals are in anti-phase, and the bottom-right plot shows the corresponding minimum correlation. In our specific application, we implement *Auto-Correlation* (where $F=G$) on the histogram of Alice’s and Bob’s pulse arrival times. This technique aims to automatically identify the relative temporal delay between the two quantum nodes by locating a distinct peak in the autocorrelation function, excluding the central peak at zero delay.

conditions. The primary challenge for the proposed solution is to distinguish the signal peaks, which represent the relative delay, from the background noise.

Figure [5.23] represents the autocorrelation result from a scenario with a low signal-to-noise ratio (SNR). We can note that in this plot, only a single, dominant peak at a lag of 0 is visible. This central peak confirms that each pulse correlates strongly with itself. However, due to the high level of background noise and a wide time window, the smaller peaks that correspond to the correlation between Alice’s and

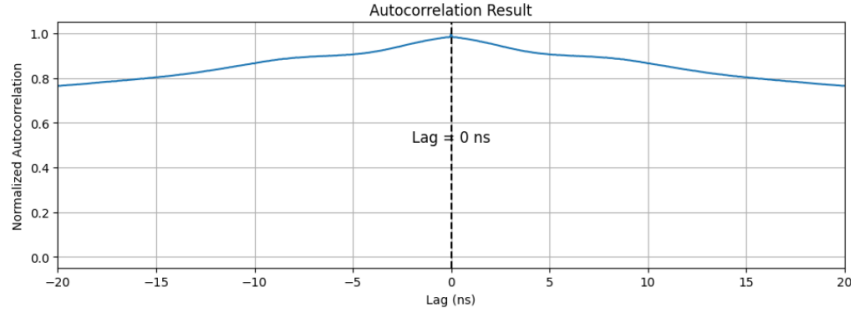


Figure 5.23: Autocorrelation result from a non-optimized, high-noise scenario. The high background noise obscures the smaller peaks, making it impossible to automatically identify the temporal delay.

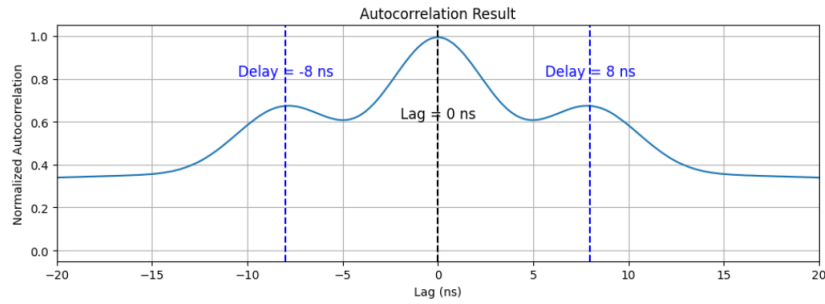


Figure 5.24: Autocorrelation result from an optimized, low-noise scenario. The two distinct side peaks at ± 8 ns clearly indicate the precise temporal delay between the pulses from Alice and Bob.

Bob's pulses are entirely obscured. This result clearly demonstrates why manually identifying the delay from a raw histogram is unreliable in a noisy environment.

In contrast, Figure [5.24] presents the autocorrelation result after the experimental parameters were optimized to reduce background noise. We can observe that this plot reveals three distinct, well-defined peaks. The central peak at a lag of 0 ns remains, verifying the autocorrelation of each pulse with itself. Crucially, two new, smaller peaks appear symmetrically at lags of ± 8 ns. These side peaks are the key result, as they represent the cross-correlation between the Alice and Bob pulses, and from their position on the x-axis, it's possible to derive a precise, objective, and automated measurement of the temporal delay between their two signals.

In conclusion, this optimized result proves that the autocorrelation algorithm, when applied to a clean signal, offers a reliable and repeatable method for measuring the temporal delay. The measured value then serves as feedback for Charlie's FPGA, which uses it to adjust the coarse and fine delays to compensate for the temporal

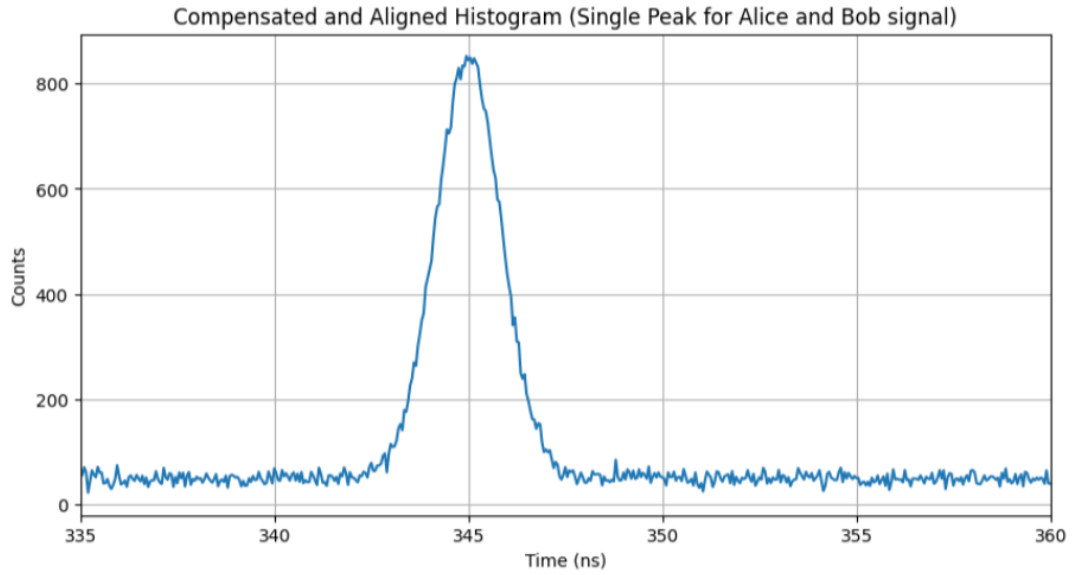


Figure 5.25: Result of the temporal alignment. This histogram shows the final, single peak after successfully compensating for the relative delay between Alice’s and Bob’s signals.

misalignment, verifying its success in real time. An example of successful alignment is shown in the histogram reported in Fig.[5.25].

5.6 Building a QKD System: Components and Their Characterization

The successful implementation of any QKD protocol requires a thorough characterization of all its components, from the laser sources and modulators to the communication channels and detectors. This characterization is fundamental for understanding the setup’s behavior and for identifying potential issues that could compromise the experiment’s success or the protocol’s security. These factors directly impact both the final QBER and the SKR of our simulations (Section (2.2.3))

Ultimately, these characterizations must be fully known to Alice and Bob before they can securely exchange a private key via QKD (Fig.[3.3]).

NOTE It is crucial that Alice’s and Bob’s experimental setups are as symmetric as possible in terms of hardware characteristics. This symmetry (understood as a balance in component properties, such as the ER of the modulators or detector

efficiencies) is a fundamental assumption for the protocol’s security, as any imbalance could be exploited by an eavesdropper to gain information.

The following section is dedicated to a detailed discussion of the technologies used in our setup, justifying our component choices and characterizing their behavior.

5.6.1 Sources



Figure 5.26: The two types of lasers used in our *setup*: a narrow-linewidth fiber laser manufactured by NKT as a sensing laser, and a diode laser manufactured by RIO as a quantum laser.

In our setup, two types of lasers were used for each terminal node of the network (Alice and Bob): a narrow-linewidth fiber laser manufactured by NKT as a sensing laser, and a diode laser manufactured by RIO as a quantum laser [Fig. 5.26]. The choice of this hybrid configuration is based on the specific properties and significant technical differences between the two types of sources.

In general, for the purposes of our QKD experiment, their main distinction lies in their frequency stability, which is determined by their internal operating mechanism and resulting linewidth. Let us investigate this choice in more detail:

- The **sensing signal** requires exceptional phase stability because it is used directly to cancel the fast phase fluctuation. For this purpose, we selected a narrow-linewidth laser, which is characterized by a very narrow linewidth of approximately 100 Hz.

Using a laser with low intrinsic noise is crucial. In fact, the phase stabilization system—a PID (proportional-integral-derivative) loop—is designed to reduce noise from both the laser and the optical fibers. Since the closed-loop

noise is the open-loop noise divided by the PID gain, a low-noise laser ensures that even with a limited-bandwidth actuator, the residual phase noise remains minimal. This is essential for the protocol's stability and performance.

- **The quantum signals**, also require low intrinsic phase noise, to ensure their coherence in the broadest possible frequency range. To ensure this, we phase-lock them to the sensing laser. PLL ensures in turn that the two sensing (and hence, the two quantum lasers) are coherent. Consequently, the quantum laser inherits the stability and narrow linewidth of the sensing laser. For this reason, we determined that less stable diode lasers, despite their wider intrinsic linewidth, are sufficient for this experiment.

It's important to note that this was a specific choice based on the available equipment. There is nothing in the protocol that would prevent the use of a narrow-linewidth laser, such as the NKT, for the quantum signals as well.

Let us now characterize the two lasers employed in our experiment. The power

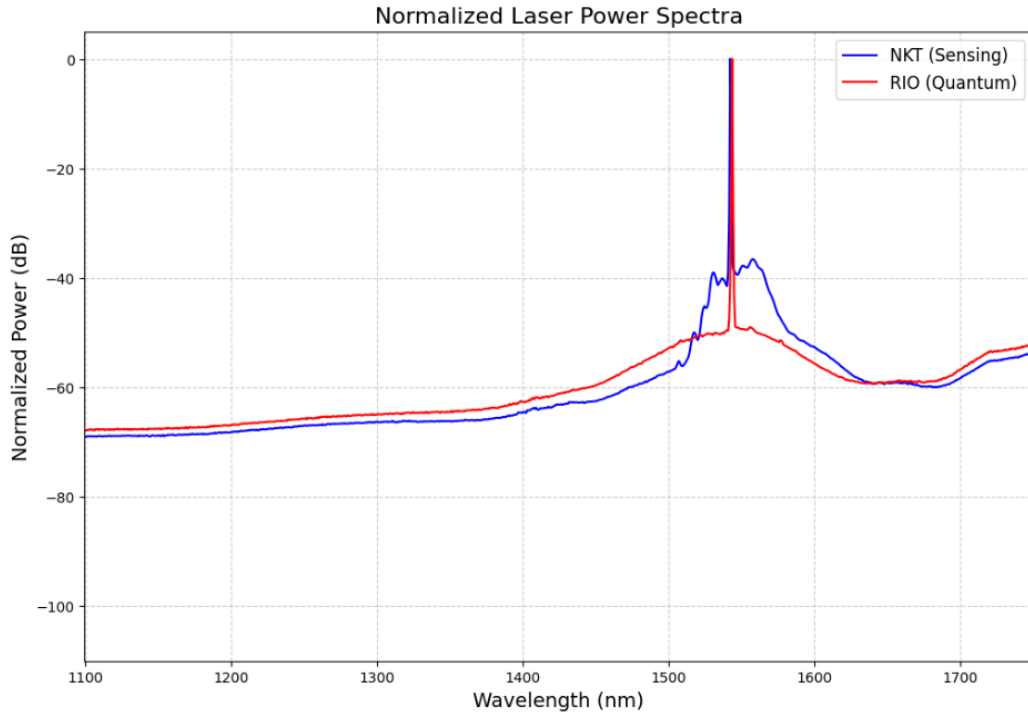


Figure 5.27: The power spectra as a function of wavelength are shown for the NKT (1542 nm) and RIO (1543.3 nm) lasers.

spectra for the NKT sensing laser (blue line) and the RIO quantum laser (red line)

are shown in Fig.[5.27]. An analysis of these spectra is fundamental to explain the specific wavelength choice in our QKD experiment: the quantum RIO laser is set at 1543.3 nm, while the sensing NKT laser is set at 1542 nm.

The primary reason for this wavelength separation is the need to keep the two lasers sufficiently close to each other without having them directly overlap. The two wavelengths must be close to ensure that the sensing and quantum signals experience similar environmental noise, such as thermal fluctuations in the fiber. This similarity ensures that the phase corrections applied to the quantum signal, which are based on noise measurements of the sensing signal, are both accurate and effective.

Conversely, it is crucial that the two lasers are not at the exact same wavelength for several reasons. First, the signals must be able to be multiplexed onto the same optical fiber using two distinct channels in a DWDM architecture. Second, a difference in wavelength is necessary for the fast phase noise compensation system (Section 5.3). This is required because we must demultiplex the two lasers and specifically utilize the phase drift information carried by the sensing laser to correct the relative phase difference between the twin field.

The noise performance of the laser sources is a critical factor for ensuring the phase coherence required by the protocol. As shown in Fig.[5.28] comparing the two sources, the phase noise of the NKT laser is significantly lower than that of the RIO laser across the entire frequency spectrum. Both lasers, however, exhibit the characteristic $1/f$ noise at low frequencies, which is a dominant noise source in these components.

The NKT laser's superior noise performance validates its classification as a narrow-linewidth laser, making it a suitable commercial alternative to the complex and expensive ultra-stable laser sources typically used for these applications. The use of this specific laser, therefore, provides a clear advantage for the practical implementation of the protocol (5.3.1).

5.6.2 Detectors

Introduction to Single Photon Detectors

To implement any discrete variable QKD (DV-QKD) protocol, it is necessary to use specific detectors capable of sensing signals down to the single-photon level. These devices are generically called *Single Photon Detectors* (SPDs), and today, a variety of technologies exist for their realization, which in turn determine their cost, performance, and scalability.

In this context, a crucial parameter for evaluating a detector's performance in

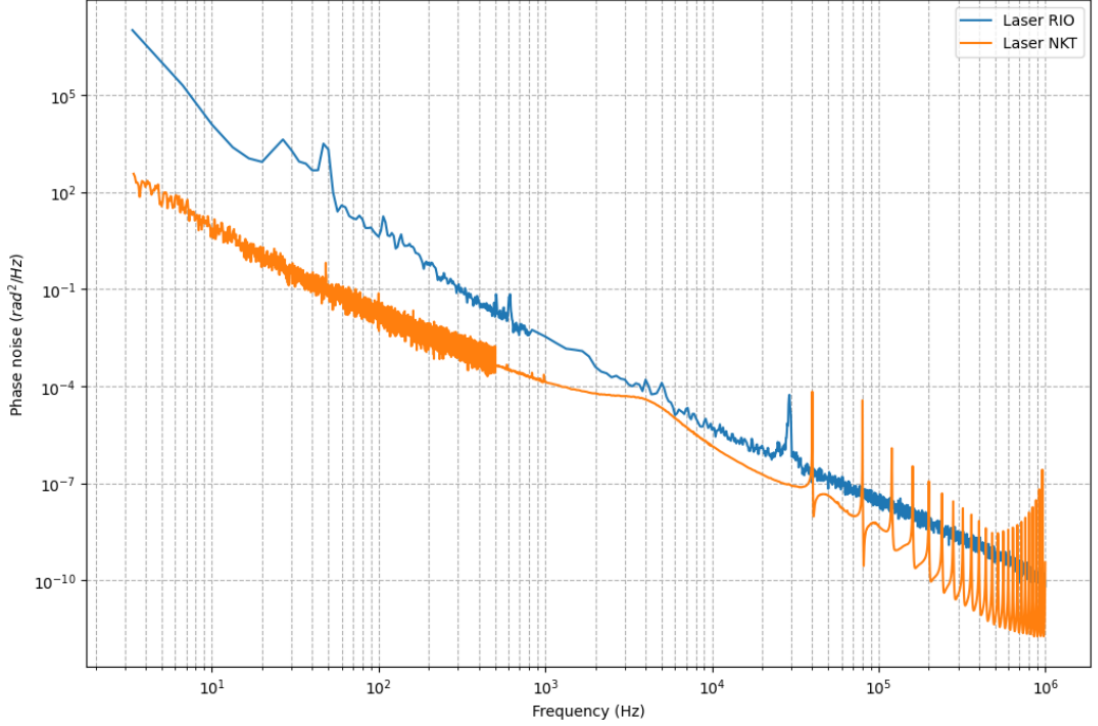


Figure 5.28: Comparison of the phase noise spectra for the RIO and NKT lasers. This figure highlights how the phase noise of the NKT laser is significantly lower than that of the RIO laser, confirming its suitability as a narrow-linewidth source for the TF-QKD protocol.

quantum cryptography is the *Dark Count rate* (P_{DC}) [51]. It represents a main limiting factor for the maximum communication distance in a QKD system, as it is directly related to the intrinsic noise of the detector itself—its tendency to register noise counts in the absence of an optical signal. More precisely, the dark counts per transmitted pulse, p_{DC} , are given by the ratio between the dark count rate P_{DC} and the system clock frequency ($p_{DC} = P_{DC}/\nu_s$). As previously discussed (Section 2.2.3), this parameter is a decisive contributor to the performance of a generic DV-QKD protocol, and particularly to the total QBER.

A more detailed technological analysis and comparison of the typical properties of APDs and SNSPDs—the two most common solutions for QKD protocols—can be found in the Appendix E

Detectors for our Experiment

For our experimental TF-QKD setup, we utilized two SPADs (D_0 and D_1) as single-photon detectors in the central node, Charlie. These detectors are responsible

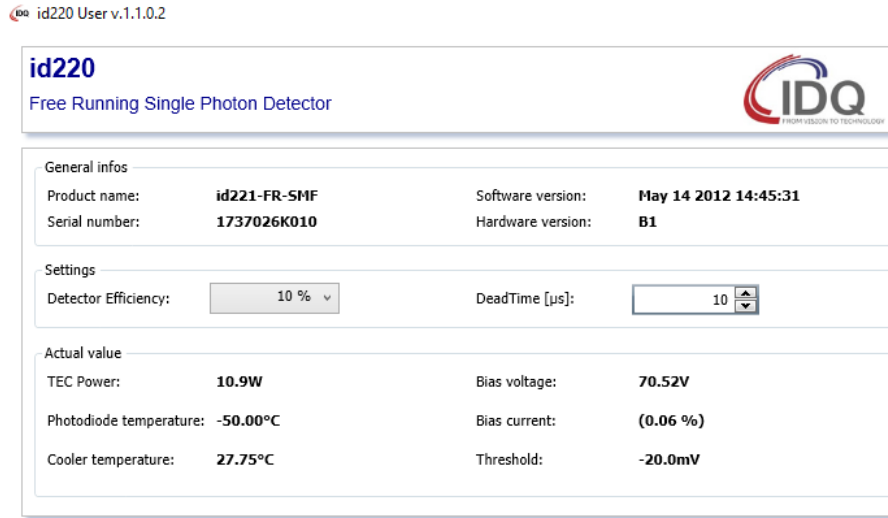


Figure 5.29: Setup and operating parameters of the ID220 SPAD detector. The graph shows the user interface of the single-photon detector (SPAD), highlighting the configuration parameters used in our setup: a detection efficiency of 10% and a dead time of 10 μ s, essential values for system characterization and power budget.

for recording the result of the quantum interference between the signals sent by Alice and Bob. The real-world performance of these devices is a critical factor that directly contributes to the total QBER of the protocol (Section 2.2.3).

The SPAD technology is based on the principle of a single photon triggering a cascade of charge carriers, a phenomenon known as an avalanche. To ensure the detector is ready for a new event, this avalanche is rapidly quenched by a dedicated circuit, which introduces a *dead time* during which the detector is unresponsive [51]. The consequences of this operating principle are crucial for understanding the limitations of the protocol. The following aspects are of particular relevance to our analysis:

- **Afterpulsing:** This phenomenon occurs when charges trapped in the semiconductor after an avalanche is quenched are released later, triggering a new, unwanted avalanche. This spurious signal is a source of noise that contributes to the QBER.
- **Dark Counts:** An avalanche can also be triggered spontaneously by a charge carrier not originating from a photon. This form of intrinsic noise directly contributes to the QBER of the protocol.
- **Efficiency:** Due to the phenomena described above and the detector's dead time, the probability that an incoming photon is correctly detected is never

100%. In our analysis, we used an efficiency of 10%.

NOTE Our setup employed two ID Quantique SPADs configured with a detection efficiency of 10% and a dead time of 10 μs (Fig.[5.29]).

Detector Background Noise

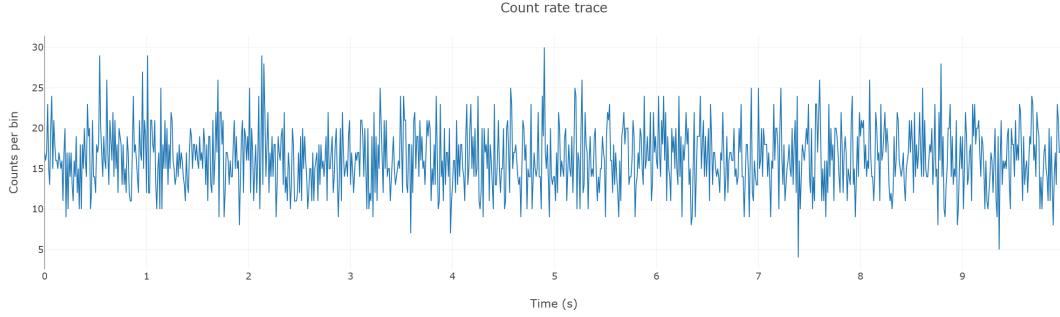


Figure 5.30: A trace of the detector’s background count rate over time, measured in counts per second, is shown in the graph. The graph illustrates the intrinsic background noise and dark photons, showing an average value of approximately 20 counts with an integration window of 100 ms and a standard deviation of about ± 10 counts. This measurement was performed using the implemented `Count_Counter` class.

To investigate the background noise that characterizes the chosen SPADs (D0 and D1) in Charlie, we first note that the distribution of incident photons is Poissonian. This implies that the standard deviation (σ) of the counts is proportional to the square root of the mean (λ), following the relationship $\sigma = \sqrt{\lambda}$.

Consequently, to estimate the detector’s background noise, we utilized the previously developed `Count_Counter` class (Section 5.12). As shown in Fig.[5.30], a measurement interval of 100 ms produced an average value of approximately 20 counts, with a standard deviation of about ± 10 counts.

This value represents the intrinsic background noise of the detectors, which is a crucial contribution to the total QBER. Accounting for this real-world detector behavior is essential for the security analysis of the SNS TF-QKD protocol (Section 3.5.3).

5.6.3 Raman Noise Analysis

In a system like ours, which implements dual-band noise detection and cancellation strategies, the primary source of background photons in the quantum channel is the

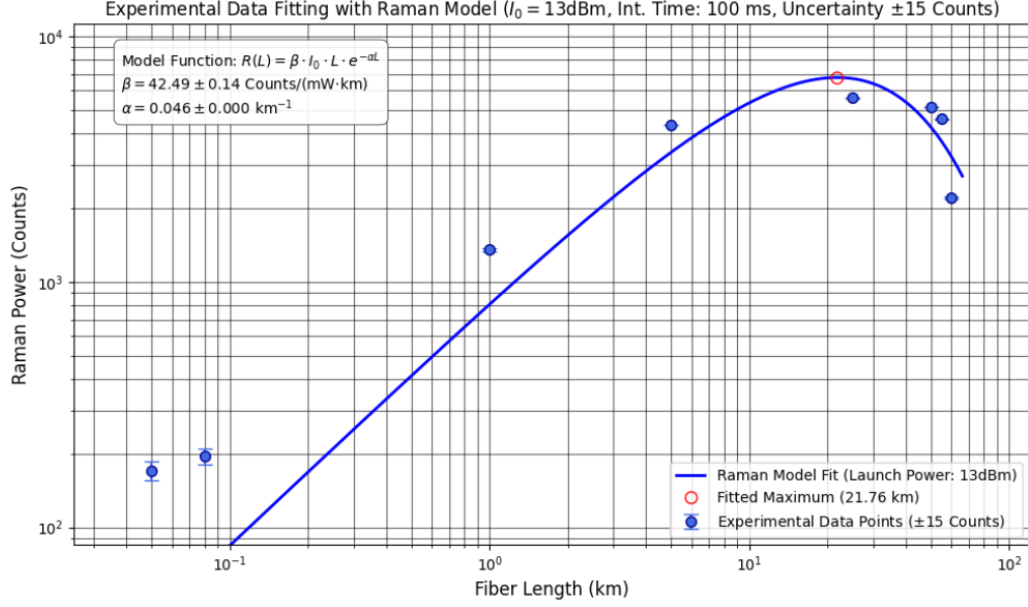


Figure 5.31: Experimental data fitting with the Raman model. The graph shows the measured Raman power (blue points) as a function of fiber length for a fixed sensing laser launch power ($I_0 = 13\text{ dBm}$). A non-linear fitting algorithm was used to apply the theoretical model, $R(L) = \beta \cdot I_0 \cdot L \cdot e^{-\alpha L}$, to the experimental data. The best-fit curve (solid blue line) and the corresponding parameters β and α are displayed. The red circle indicates the maximum *Effective Length* at which the Raman noise power reaches its highest value before beginning to decrease due to fiber attenuation.

spontaneous Raman effect. This is an inelastic, broadband scattering phenomenon induced by photon-phonon interactions within the fiber [12]. These Raman photons are generated by the sensing laser, which travels in the same fiber as the QKD laser. A fraction of these generated photons falls into the band of the QKD laser and thus represents the main source of background noise in our experiment.

The analysis of background noise in our setup is based on a mathematical model that describes the power of the Raman signal (R) as a function of the fiber length (L). The derivation of this model is based on a differential analysis of the Raman signal power generated in a small segment of fiber, followed by integration over the entire fiber length. Let's analyze this mathematical derivation in more detail.

The mathematical model for the measured Raman power is derived by considering the generation and subsequent attenuation of the Raman signal along the optical fiber.

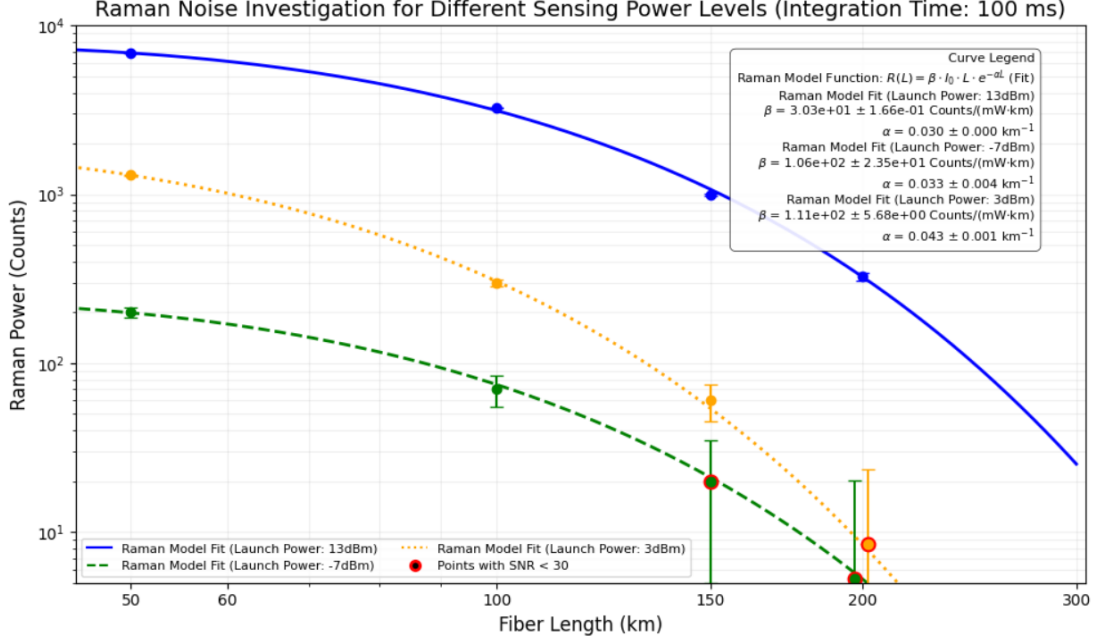


Figure 5.32: Raman noise as a function of fiber length for three different sensing laser launch powers: 13 dBm, 3 dBm, and -7 dBm. Each curve illustrates how the Raman noise scales with distance for a specific launch power. The red circles indicate the lengths at which the signal-to-noise ratio (SNR) of the beat note drops below a critical threshold of 30 dB. The graph demonstrates that even at low launch powers, an acceptable SNR can be maintained for long distances, suggesting that a total communication range of up to 200 km is feasible for the TF-QKD protocol in this setup.

At any point x along the fiber, the power of the pump signal decays exponentially due to attenuation. The intensity of the light $I(x)$ at a distance x from the source is given by:

$$I(x) = I_0 \cdot e^{-\alpha x} \quad (5.38)$$

where I_0 is the initial light intensity at $x = 0$, and α is the fiber's attenuation coefficient.

Consider a small, infinitesimal segment of fiber with length dx located at position x . The Raman signal generated in this segment (dR) is proportional to the intensity of the pump light at that point and the length of the segment itself. This proportionality is defined by the Raman generation coefficient, β :

$$dR = \beta \cdot I(x) \cdot dx \quad (5.39)$$

Note: The coefficient β is a function of wavelength. We do not make this dependence explicit for the sake of clarity.

Substituting the expression for $I(x)$ from Equation 5.38, we get the differential generation equation:

$$dR = \beta \cdot I_0 \cdot e^{-\alpha x} \cdot dx \quad (5.40)$$

This differential equation describes the amount of Raman signal generated in each infinitesimal segment of the fiber.

The Raman signal generated in the segment dx at position x must, in turn, travel the remaining length of the fiber ($L - x$) to reach the detector located at the end L . This signal also undergoes attenuation. For simplicity, it is assumed that the attenuation coefficient is the same for both the Raman signal and the pump signal, i.e., α . The power of the Raman signal (dR_f) that arrives at the detector, after being attenuated over the path $L - x$, is given by:

$$dR_f(\text{at detector}) = dR \cdot e^{-\alpha(L-x)} \quad (5.41)$$

Substituting the expression for dR from Equation 5.40:

$$dR_f = (\beta \cdot I_0 \cdot e^{-\alpha x} \cdot dx) \cdot e^{-\alpha(L-x)} \quad (5.42)$$

Simplifying the exponents:

$$dR_f = \beta \cdot I_0 \cdot e^{-\alpha x - \alpha L + \alpha x} \cdot dx \quad (5.43)$$

The terms involving x cancel out, leading to:

$$dR_f = \beta \cdot I_0 \cdot e^{-\alpha L} \cdot dx \quad (5.44)$$

This result shows that the contribution of each individual fiber segment to the total measured signal at the detector is a constant, $\beta \cdot I_0 \cdot e^{-\alpha L}$, multiplied by the segment length dx .

To obtain the total Raman power (R) measured at the detector, we must sum the contributions of all the elementary segments of the fiber, from the initial point $x = 0$ to the final point $x = L$. This is achieved by integrating the expression for dR_f over the entire length of the fiber:

$$R = \int_0^L dR_f \quad (5.45)$$

$$R = \int_0^L (\beta \cdot I_0 \cdot e^{-\alpha L}) \cdot dx \quad (5.46)$$

Since the terms β , I_0 , and $e^{-\alpha L}$ are constant with respect to the integration variable x , they can be taken outside the integral:

$$R = \beta \cdot I_0 \cdot e^{-\alpha L} \int_0^L dx \quad (5.47)$$

The integral of dx from 0 to L is simply L .

$$R = \beta \cdot I_0 \cdot L \cdot e^{-\alpha L} \quad (5.48)$$

This is the final model used for fitting our experimental data.

NOTE This equation reflects the balance between two fundamental physical phenomena: the Raman photon generation rate, which decreases along the fiber, and signal attenuation. In particular, from Equation 5.48 we can note that the former is dominant at short distances from the source, as the pump light intensity is strongest, while for long distances, the attenuation dominates the total power budget.

Experimental Analysis of the Raman Noise Model

Fig.[5.31] illustrates the trend of Raman noise as a function of fiber length for a fixed launch power of 13 dBm. The measurements were conducted using a low-noise InGaAs/InP avalanche photodiode with a quantum efficiency of 10 % and a dead time of 10 μs . Under these conditions, we recall that the detector's intrinsic background photon rate was about 20 counts (with an uncertainty of ± 10 counts) in 100 ms (Section 5.6.2).

The most significant result of the analysis is the existence of a characteristic length at which the Raman noise reaches its maximum. This length, indicated by a red circle in the graph, is about 20 km. At shorter distances, the Raman photon generation dominates over the fiber attenuation. Conversely, at longer distances, the attenuation is dominant. From a practical point of view for a QKD protocol, it is advantageous to operate with fiber lengths greater than this value, as the contribution of Raman noise tends to decrease, consequently reducing the QBER.

It is also interesting to note how the first experimental values deviate significantly from the fitting curve. This deviation occurs because at short distances, the predominant background noise is not of Raman origin, but stems from the intrinsic noise of the detector. The contribution of Raman noise becomes relevant only for distances greater than about 5 km, where it begins to grow in an almost linear manner until it reaches the highlighted maximum peak.

As a second test, we quantified the impact of the pump light intensity on Raman noise

scales linearly with the pump light intensity. While reducing the power of the sensing laser would clearly reduce the impact of Raman noise, we must compromise this reduction with the need to ensure an adequate SNR for the beatnote used in the stabilization process.

In this measurement, we consider a pair of sensing lasers transmitted from Alice and Bob to Charlie. Simultaneously to the Raman count measurement, we record the SNR of the beatnote between the sensing lasers at Charlie’s node. We, therefore, repeated the previous measurement with different launch powers as a function of fiber length. This analysis aims to determine the maximum usable fiber length for a given launch power before the SNR drops below a critical threshold. Exceeding this threshold would result in an excessive number of Raman background photons, leading to an unacceptably high QBER.

Fig. [5.32] illustrates the Raman noise as a function of fiber length for three different launch powers: 13 dBm, 3 dBm, and -7 dBm. For each curve, the critical points (red circles) indicate the lengths at which the beatnote’s SNR drops below the threshold of 30 dB.

The reduction in launch power leads to a clear decrease in the generated Raman noise, as shown by the downward progression of the curves. The most significant result is that, even with a very low launch power of -7 dBm (0.2 mW), it is possible to maintain an acceptable SNR for the beatnote even for lengths on the order of 100 km. This result suggests the possibility of implementing TF-QKD communication over a total distance of 200 km in a real-world setting, with each Alice-Charlie and Bob-Charlie link extending 100 km. At this distance, the Raman noise is on the order of 50 counts in 100 ms, a manageable contribution that, considering the intrinsic noise of the detector used simultaneously (on the order of 20 counts in 100 ms), does not compromise the performance of the protocol.

5.6.4 Encoder and Signal Modulation

The purpose of the encoder is to take a continuous-wave quantum laser signal as input and allow Alice (or Bob) to modulate it in both intensity and phase to implement the SNS TF-QKD protocol with decoy states. Specifically, the encoder receives the continuous signal and converts it into a pulse train with a frequency of 250 MHz, setting the intensity and phase of each pulse i according to the protocol requirements. The optical components are housed in a protective enclosure to avoid thermal drift of the modulators’ bias point.

An example of modulation pattern is showed in Fig.[5.33]. Here is important to note that within each a cycle of duration 4 ns, the modulated pulse has a duration of 2 ns and it is followed by an empty buffer of 2 ns. This buffer prevents signal

contamination between consecutive pulses and provides a necessary margin for phase modulation accuracy.

In our experiment, the chosen TF-QKD protocol requires two types of modulation

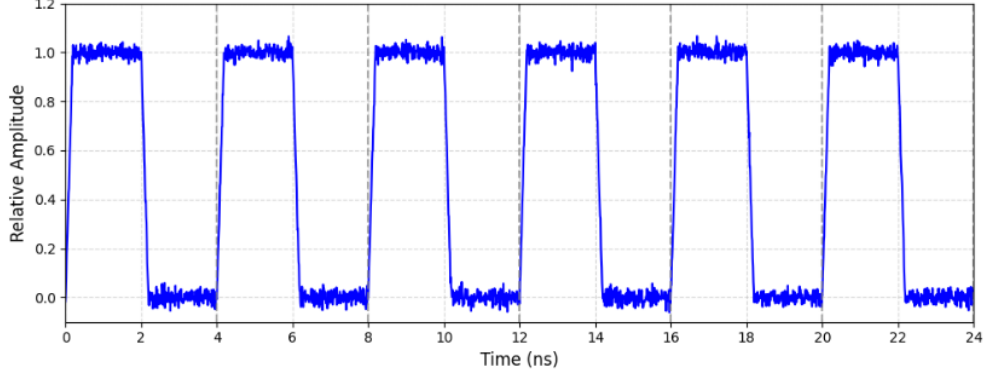


Figure 5.33: A simplified representation of the encoder's modulation pattern. The signal features a 2 ns pulse followed by a 2 ns empty buffer, resulting in a total period of 4 ns.

that our encoder, based on a series of cascaded modulators (presented in Fig. [5.7]), must perform:

- By combining the action of the two intensity modulators (IM1,2), pulses of five different intensity levels are prepared: μ_r (*signal for realignment*), μ_Z (*signal state or send state*), μ_2 (*strong decoy*), μ_1 (*medium decoy*), and μ_0 (*not-send state and weak decoy = vacuum state*).
- The PM encodes the phase of each quantum signal pulse with one of 16 phase values ($\theta \in \{0, \pi/8, 2\pi/8, \dots, 15\pi/8\}$), satisfying the requirements for phase randomization and qubit encoding.

NOTE As describe in Section 5.4, despite our setup being based on the dual-band technique that generally allows for long duty cycles ($d > 0.9$), the phase compensation process still requires the QKD phase to be periodically interrupted. For this reason, a pulsed signal with intensity (μ_r), higher than the decoy states and not phase-modulated, is sent from Alice and Bob to Charlie.

The design of the encoder's modulation stage, particularly the use of two cascaded intensity modulators, is driven by a critical system requirement: the ER must be equal to the total optical link attenuation (Section 5.6.4). In particular, this necessity derives from two contradictory conditions the protocol must satisfy across distinct temporal windows:

- **Realignment Window:** A minimum of one photon per pulse must reach Charlie's central node for accurate synchronization.
- **QKD Window:** The "signal state" must be attenuated to a single-photon level for secure key exchange, meaning the number of photons per pulse arriving in Charlie is equivalent to the link transitivity.

The long-distance nature of TF-QKD necessitates considering link distances on the order of 200 km (as a minimum benchmark). Consequently, assuming a standard fiber attenuation coefficient of approximately 0.2 dB/km, this determines an optical link loss of up to 40 dB. As previously shown, this loss directly imposes a consequential ER requirement of 40 dB. If the quantum laser power were set strictly for the single-photon QKD condition, the signal would be too attenuated for the essential realignment procedures.

To overcome this, a much higher initial power (P'_{in}) must be launched during the Realignment Window to ensure sufficient photon count at Charlie. This high-power logical '1' state, designed to overcome the 40 dB link loss, must be attenuated down to the single-photon level during the QKD Window, while the logical '0' state must approach zero photons. This required power difference between the '1' and '0' states directly necessitates a minimum ER of 40 dB. Since commercial single-stage modulators typically provide an ER of only around 20 dB, the setup must employ a cascaded configuration with two modulators, each with an ER of at least 20 dB, to collectively meet the critical 40 dB requirement. The characterization of the four intensity modulators is reported in detail in Section (5.6.4).

Ultimately, we need to consider that the system's real characteristics further complicate the power budget. In fact, the final power available for interference is reduced not only by the fiber attenuation but also by internal losses. Characterization revealed a loss of 26.4 dB at Alice's node and 27.1 dB at Bob's node and an additional attenuation of 14.5 dB within the central Charlie node itself. Therefore, accurate estimation of the input power during the Realignment Window is crucial to ensure that, even after all these losses, at least one photon per pulse still arrives at Charlie's node, thereby satisfying the protocol's requirements. The characterization of the single-photon regime for phase realignment is reported in Section 5.6.6.

Characterization of Intensity Modulators

The behavior of an intensity modulator [52] is governed by three fundamental voltage parameters: the DC bias voltage (V_o), the applied modulation signal voltage (V_m), and the half-wave voltage (V_π). In particular, the DC bias voltage (V_o) sets the initial operating state and the output optical power level; the half-wave voltage

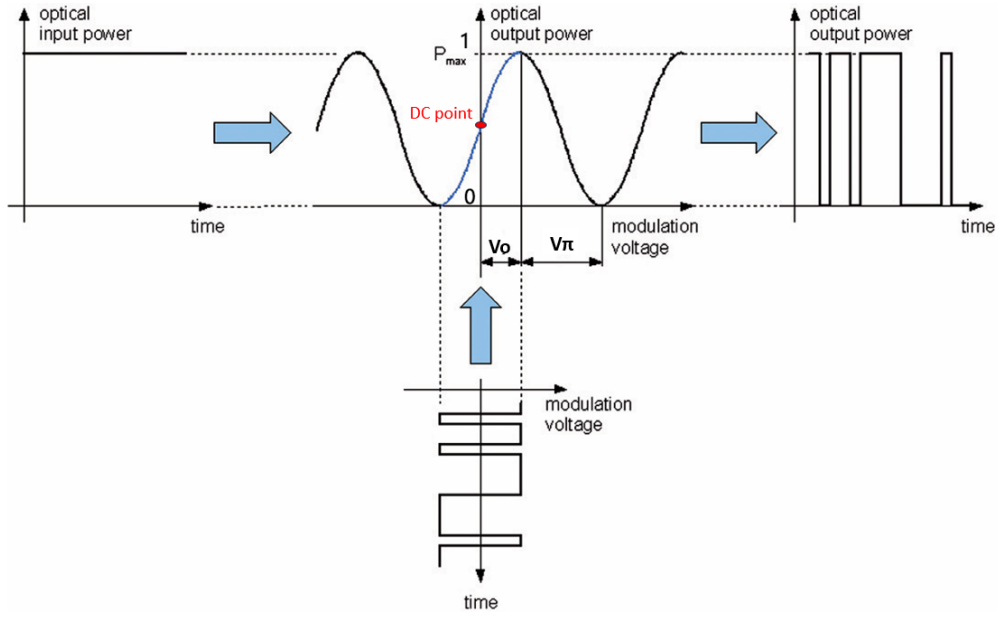


Figure 5.34: Standard Modulator Characteristic [52]: A typical modulator operation where V_0 is set at the point of maximum slope for linear modulation. A DC bias point is applied, and an RF signal modulates around it.

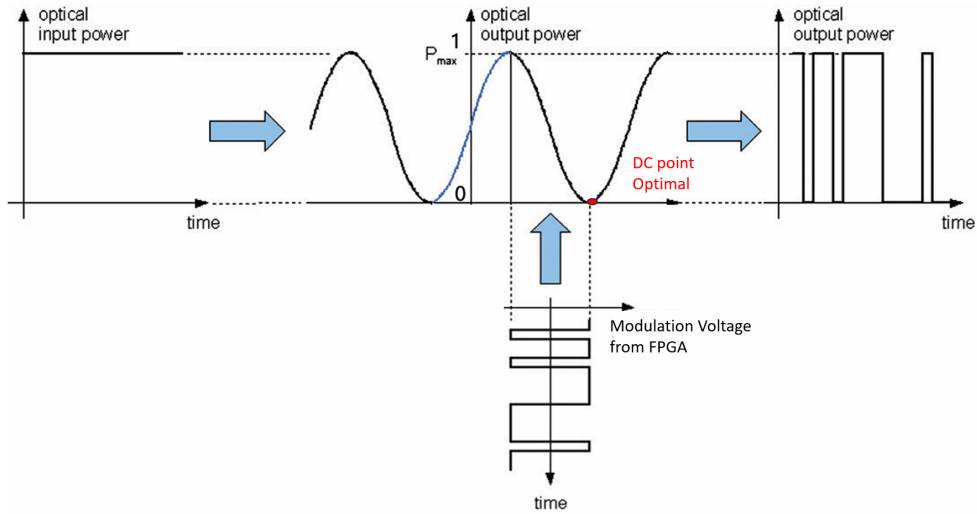


Figure 5.35: Our Modulator Characteristic with Inverted Logic: Our specific setup is calibrated such that the optimal DC point corresponds to minimum optical output. A negative voltage from the FPGA generates a high optical pulse (logical "1"), while 0V maintains minimum output (logical "0").

(V_π) is a characteristic parameter of the modulator, specifying the voltage change required to transition the optical transmission from its minimum to maximum; and the modulation signal voltage (V_m) is the signal applied to the modulator to dynamically change the intensity.

In a standard Mach-Zehnder configuration, the transmission curve shows a sinusoidal response, as illustrated in Fig. [5.34]. For linear modulation, the bias point is typically set to the region of greatest slope, usually at $V_o = V_\pi/2$, where a positive modulation signal determines an increase in the output optical power. Crucially, to utilize the full dynamic range of the modulator, the applied modulation signal voltage must satisfy $V_m = V_\pi$. Only this condition ensures that the intensity can be switched between its maximum and minimum transmission points.

However, our specific setup employs a Digital-to-analog converter (DAC) and modulators with a unique configuration that deviates from this conventional relationship. In general, the DAC driving the modulator is configured to produce both a signal and its negation, however our system does not utilize the full differential signal but only one of the two output lines. This design choice results in a non-symmetric voltage range, leading to the following behavior, as illustrated in Fig.[5.35]: our system operates such that a more negative voltage signal from the FPGA corresponds to a higher optical intensity (a "logical 1"), while a more positive voltage (specifically, 0V in our case) yields a minimum optical intensity (a "logical 0"). This behavior is directly linked to the hexadecimal voltage values provided as 16-bit inputs to the FPGA. The FPGA utilizes a two's complement representation, where `0x8000` is the most negative value and `0x7FFF` is the most positive.

For this setup, it is necessary to manually determine the optimal DC bias point (V_o) at the beginning of each experimental session, as indicated, for instance, by the "DC point Optimal" marker in Figure [5.34]. This voltage intrinsically corresponds to the minimum transmission state, meaning that with no active modulation voltage (i.e., maintaining 0V), the modulator produces negligible optical output (our "logical 0", or P_{min}). This effectively simplifies the modulation process, as the "zero" state is inherently set by the bias.

To generate a "logical 1" (maximal optical power, P_{max}), the FPGA, when set to its most negative value ('`0x8000`'), applies a specific negative voltage (approximately -5 V for our modulators). This voltage drives the modulator from its optimal DC point (0 V) to the peak of its transmission curve, resulting in maximum optical output. Conversely, to generate a "logical 0", the FPGA applies a value (e.g., '`0x17FF`') that maintains the voltage at 0 V, thereby keeping the modulator at its minimum transmission state. The lower part of Figure 5.35 shows these discrete voltage levels applied by the FPGA and their corresponding optical output pulses.

This specific calibration and the strategic choice of the optimal DC point are essential for the intensity coding of quantum signals. It ensures a clear, repeatable, and robust distinction between the logical "0" and "1" states sent from Alice and Bob to the central Charlie node, despite the inverted voltage-power relationship.

FPGA Codes for Intensity and Phase Modulation

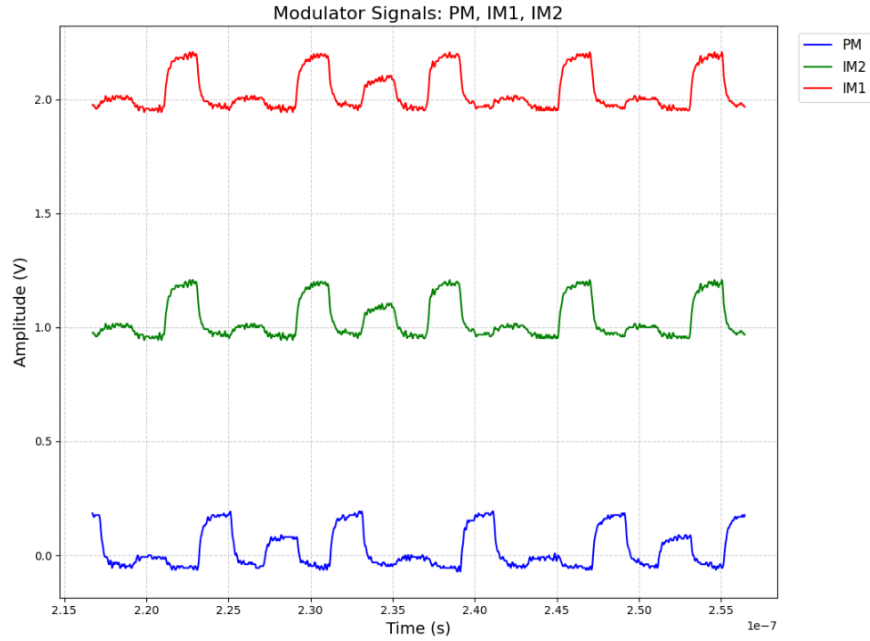


Figure 5.36: Modulation patterns generated by the twin FPGA boards for the intensity modulators (IM1 and IM2) and the phase modulator (PM). The identical intensity patterns for IM1 and IM2 (red and green lines) were implemented for a simulation of the asymmetric 4-decoy SNS TF-QKD protocol, displaying four distinct intensity levels: signal, vacuum=decoy0, decoy1 and decoy2 (Section 3.5.3). The phase modulator (PM) pattern (blue line) shows an example of a subset of the 16 distinct phase levels used for coordinated phase randomization (Section 3.2)

Based on this characterization, it is now possible to demonstrate how the intensity and phase modulation are controlled using the twin FPGAs in Alice and Bob. This will highlight the strategic importance of the implemented *alice_bob_modulator* code. This code, executed on identical FPGA boards assigned to Alice and Bob's nodes, is essential for generating specific intensity and phase modulation patterns (Fig.[5.36]). These procedures are crucial both for characterizing our

setup and for implementing the first simulations of the asymmetric 4-decoy SNS TF-QKD protocol. Ultimately, the flexibility offered by the FPGAs allows for the identification of optimal voltage levels for the modulators and the customization of modulation patterns according to the protocol's needs.

Let us now describe in more detail how the intensity and phase modulation are managed using this implemented code.

Intensity Modulation and Voltage Levels Intensity modulation is the foundational process for creating the optical signals required by the protocol. Our setup operates in a *reversed* manner: a more negative voltage signal corresponds to a larger peak in the optical pulse (representing a "one" in the modulation), whereas a positive signal generates a minimum (representing a "zero"). This behavior is intrinsic to the calibration of both the modulators and the Digital-to-Analog Converter (DAC) of the FPGAs.

Furthermore, the hexadecimal voltage values, provided as a 16-bit input, directly determine the amplitude of the pulse. Specifically, the FPGAs utilize a two's complement representation, where 0x8000 corresponds to the most negative value and 0x7FFF to the most positive. This capability allows for the generation of a wide range of intensity levels, which have been precisely mapped to specific voltage levels through laboratory calibration. In turn, these voltage levels correspond to an appropriate number of photons per pulse. Based on these values, we can select the intensity levels for the *signal*, *decoy*, and *no-signal* (vacuum) states provided by the protocol.

Phase Modulation and 16-Level Quantization For phase modulation, an optimal protocol implementation requires a discretization into $M_{opt} = 16$ levels, symmetrically distributed between 0 and 2π .

Similar to intensity modulation, the FPGAs generate these phase levels using eight distinct voltage values, which correspond to eight positive and eight negative phases. This symmetry is useful for *Coordinated Twin-Field Phase Randomization* (3.4.2), a crucial phase randomization technique in the security demonstrations of the TF-QKD protocol.

In conclusion, for both types of modulation, the FPGA code allows for loading the desired intensity and phase patterns into memory. The iterative process of writing to the MMIO (Memory-Mapped I/O) memory ensures that the FPGA board generates a sequence of light pulses that reproduces the desired pattern (Fig.[5.36]).

ER of Modulators

Moreover, a detailed characterization was conducted to quantify the performance of modulators.

For this purpose, we utilized a modified version of our code to measure the count rate on the detector (Section 5.12). This new version allowed for the automatic recording of the instantaneous count values as a function of the DC point for each modulator. From this acquired data, then the ER of each modulator was calculated.

The characterization procedure for each device was divided into three main steps to accurately quantify its ER.

1. For each intensity modulator, a data acquisition code was employed to measure the number of photons incident on the SPAD, while keeping the pulse modulation signal off. This enabled us to determine the optimal DC bias point which, as explained previously, coincides with the logical 0 in our setup.
2. Simultaneously with the count rate measurement, the modulator's DC operating point was dynamically varied to explore its entire functional range. This process enabled the identification of both the point of maximum attenuation (corresponding to the "logical 0") and the point of maximum transmission (the "logical 1").
3. From the maximum and minimum count rate values recorded during the DC point variation, the ER of each individual modulator was calculated. The ER value, expressed in decibels (dB), quantifies the ratio between the optical power transmitted in the state of maximum transmission and that in the state of maximum attenuation.

This analysis confirmed that all four modulators had an ER of at least 20 dB. The specific results are summarized in the table below.

Node	Modulator	Extinction Ratio (ER)
Alice	IM1	25.48 dB
Alice	IM2	26.68 dB
Bob	IM1	20.7 dB
Bob	IM2	20.3 dB

Table 5.1: Measured Extinction Ratios for each intensity modulator.

We can consequently conclude that the chosen modulators are suitable for the implementation of the TF-QKD protocol. With two cascaded modulators for each node, they guarantee a combined ER of at least 40 dB, as required for the correct functioning of the protocol itself.

Characterization of Amplifiers

Despite the excellent performance of the intensity modulators, it was necessary to investigate the characteristics of the RF amplifiers used to drive them with the control signal originating from the FPGA.

To do this, we generated a radio-frequency (RF) signal using a function generator and determined the input voltage that corresponded to the compression point of each amplifier. This point is defined as the state where the output voltage (V_{out}) reaches its maximum value and begins to saturate, which was observed using an oscilloscope. By repeating this procedure for all amplifiers, we empirically discovered that each component exhibited a different compression point and provided a unique output voltage for a given input voltage (e.g., 100 mV).

This inconsistency represents a significant operational problem for our experiment, especially considering the requirement for high symmetry and high extinction ratio (ER). For example, as summarized in Table 5.2, a 100 mV input yields an output of 4.8 V for Alice's Amplifier 1 (its saturation point), while the same 100 mV input provides only 3.8 V for Alice's Amplifier 2, significantly below its saturation point of 4.0 V. This inconsistency means that a nominally identical input signal cannot drive the two cascaded modulators identically, which would directly prevent the achievement of the total ER of 40 dB required by the cascaded modulator configuration.

This analysis confirmed our hypothesis that the amplifiers are characterized by different gains and saturation characteristics, providing an explanation for the unexpected modulator behavior observed during preliminary testing. Furthermore, it underscored the critical importance of individually characterizing each component in a complex experimental setup to ensure the system performs as expected.

The solution we adopted to overcome this constraint was to individually maximize the gain of each amplifier ensuring that the input control signal from the respective FPGA (where the hexadecimal value 0x8000 corresponds to the "one" state and 0x17FF to the "zero" state) was sufficient to guarantee that the ER of each individual modulator was at least 20 dB. This condition is sufficient for our experiment because the cascaded configuration of the two modulators at each node inherently guarantees the required total ER of 40 dB. This resulting ER satisfies the two critical hypotheses underlying a valid real-world implementation of the TF-QKD protocol (Section 4.1.1).

Node	Amplifier	Output Voltage (V_{out}) at -100 mV Input	Saturation Point
Alice	Amp 1	4.8 V	4.8 V
Alice	Amp 2	3.8 V	4.0 V
Bob	Amp 1	3.7 V	4.1 V
Bob	Amp 2	3.6 V	3.8 V

Table 5.2: Measured characteristics of the RF amplifiers for each node.

5.6.5 Variable Optical Attenuator

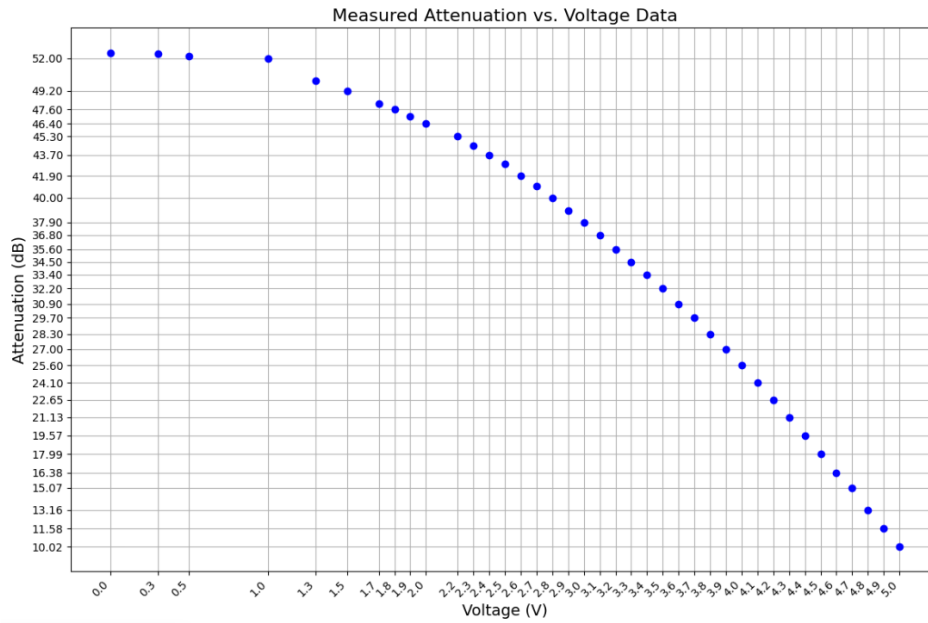


Figure 5.37: Characterization of the Variable Optical Attenuator (VOA) used in Alice and Bob’s enclosures.

To correctly implement the SNS TF-QKD protocol with decoy states, Alice and Bob’s encoder includes a Variable Optical Attenuator (VOA). This device, common in optics, varies the light’s attenuation based on a voltage signal. Specifically, in our setup, the VOA ensures an average attenuation to transition from the classical to the quantum regime and to correct for any variations in attenuation over time.

The characterization of the VOA used is shown in the graph in Fig.[5.37]. The graph illustrates that the device has a well-defined operating range, between approximately 1.5 V and 5.0 V, where the attenuation can be controlled precisely. Below 1.5 V, the VOA behaves like a fixed attenuator with a maximum attenuation greater than

50 dB, while within its operating range, the attenuation constantly decreases as the voltage increases. It is also important to note that the total attenuation range controllable by the VOA goes from approximately 50 dB to 10 dB.

5.6.6 Characterization of the Single-Photon Regime for Phase Re-alignment

This section investigates the experimental conditions under which our setup can satisfy a practical requirement of the TF-QKD protocol: **ensuring at least one photon per pulse reaches Charlie's node for temporal realignment.**

Our measurements aim to determine the precise attenuation needed to convert a high-power quantum laser signal (approximately 1 mW) into the single-photon-per-pulse regime at the detector. This is a critical condition for the successful exchange of the final secret key in our system.

The measurements were performed using a simplified setup comprising a single transmitting node (Bob) and a receiving node (Charlie). This configuration, while simplified for our measurements, is based on the assumption that Alice's and Bob's nodes and channels are subject to the same total attenuation and are symmetric in length.

To minimize background noise, particularly Raman noise, the Sensing laser was turned off, and consequently, the phase modulator as well, although its insertion losses (IL) were still considered in the power budget. For precise attenuation control, two variable optical attenuators were used in cascade before the detector: an analog VOA (with an attenuation controllable from 0 to 50 dB via external voltage, whose characterization is reported in Section 5.37) located in the Alice node, and an additional digital VOA (with a controllable attenuation from 0 to 60 dB) positioned immediately before the SPAD detector to prevent its saturation.

The experiment is based on the use of two cascaded modulators (IM1,2) that, when combined, provide an Extinction Ratio (ER) of approximately 40 dB; while the estimated input power in our system is on the order of 1 mW.

A crucial preliminary analysis is to determine the number of photons per pulse. For this purpose, we first calculate the number of photons per second for a given power. The energy of a single photon (E_{photon}) at a wavelength $\lambda = 1550 \text{ nm}$ (typical for QKD systems) is given by $E_{\text{photon}} = hc/\lambda$, where h is Planck's constant ($6.626 \times 10^{-34} \text{ J s}$) and c is the speed of light ($3 \times 10^8 \text{ m s}^{-1}$). Consequently:

$$E_{\text{photon}} = \frac{hc}{\lambda} = \frac{(6.626 \times 10^{-34} \text{ J} \cdot \text{s}) \times (3 \times 10^8 \text{ m/s})}{1550 \times 10^{-9} \text{ m}} \approx 1.28 \times 10^{-19} \text{ J} \quad (5.49)$$

In our case, given an input power (P_{in}) of 1 mW, the number of photons per second ($N_{\text{photons/s}}$) is:

$$N_{\text{photons/s}} = \frac{P_{\text{in}}}{E_{\text{photon}}} = \frac{1 \times 10^{-3} \text{ W}}{1.28 \times 10^{-19} \text{ J}} \approx 7.8 \times 10^{15} \text{ photons/s} \quad (5.50)$$

Now, since our logical '1' pulses have a duration of 2 ns, the number of photons per pulse ($N_{\text{photons/pulse}}$) was estimated as:

$$N_{\text{photons/pulse}} = (7.8 \times 10^{15} \text{ photons/s}) \times (2 \times 10^{-9} \text{ s}) \approx 1.6 \times 10^7 \text{ photons/pulse} \quad (5.51)$$

At this point, we need to appropriately calibrate the system's attenuation. The experimental procedure was structured as follows:

- We set the Bob transmitter in the condition of initial attenuation: the analog VOA was set to its maximum attenuation of 50 dB, and the digital VOA was set to a fixed 5 dB, resulting in a total initial attenuation of 55 dB.
- Initially, the optimal DC (Direct Current) operating points for the two modulators were set. Using the `COUNT_RATE` class (Section 5.12) to monitor the SPAD detector's counts, the number of photons corresponding to the logical '0' state was observed. This value, approximately 20 counts in 100 ms, represents the intrinsic background noise of the SPAD (Section 5.6.2). This condition defines the vacuum state, representing the lowest achievable attenuation.
- Subsequently, the measurement procedure involved gradually decreasing the attenuation. This was achieved by adjusting the control voltage of the analog VOA in increments of 100 mV, while the digital VOA was set to its minimum attenuation state (0 dB). The detector count was continuously monitored over an integration time of 100 ms, progressively rising as the total attenuation decreased.

Results Analysis

The count continued to rise until it reached a saturation value of 5000 clicks in a 100 ms interval. This value is considered the maximum theoretical count possible due to the detector's quenching effect, which prevents the detection of all photons for rates exceeding 50 kHz (equivalent to 5000 clicks/100 ms). To verify this, we must consider that for this specific measurement, the modulation clock was set to a frequency of 50 kHz which corresponds to an interval of 20 μ s between two consecutive logical '1' pulses. Therefore, the number of pulses expected in a 100 ms interval is:

$$N_{\text{pulses}} = \frac{100 \text{ ms}}{20 \mu\text{s}} = \frac{100 \times 10^{-3} \text{ s}}{20 \times 10^{-6} \text{ s}} = 5000 \quad (5.52)$$

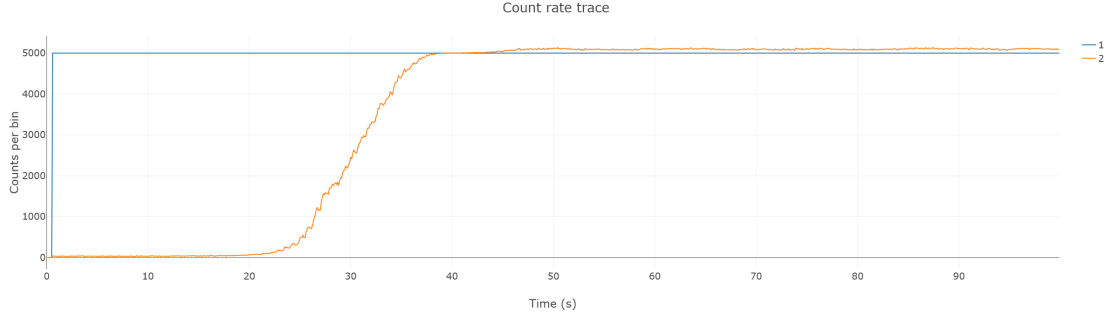


Figure 5.38: The graph shows the progressive increase of the detector count (orange line) as the voltage supplied to the analog VOA increases (corresponding to a decrease in optical attenuation). The count stabilizes at 5000 clicks per second (blue line), a value that confirms the achievement of the single-photon regime and coincides with the theoretical number of pulses emitted in a 100 ms interval.

From this simple theoretical calculation, we deduce that reaching 5000 clicks in a 100 ms interval indicates that the system is receiving at least one photon per emitted pulse. This condition corresponds exactly to the fundamental requirement of the protocol for the realignment window, namely the critical hypothesis described previously.

The fixed optical attenuation required for phase realignment was calibrated to a maximum value of 13.16 dB, achieved by setting the VOA to a control voltage of 4.80 V. This calibration is crucial as it simultaneously ensures the transition from the classical to the quantum regime while stabilizing the photon count at the maximum theoretical value of 5000 clicks (thus preventing detector saturation).

NOTE After exceeding 4.80 V supplied to the VOA, the count began to exceed the expected 5000 clicks. Obviously, according to the theoretical calculation shown, this behavior is not possible given our choice of parameters for this experiment and must therefore be interpreted as a non-ideal detector anomaly. It should be remembered that a SPAD, operating in Geiger mode [12], is not able to distinguish the number of photons per single click; consequently, even in the case of multiple photons per pulse, the count should not exceed the total number of single-photon pulses estimated. This anomaly is attributable to the high number of photons that exceed the detector’s management capacity; specifically, it is likely that the quenching circuit is not able to correctly interrupt the avalanche of charge carriers generated after the first photon, compromising the count.

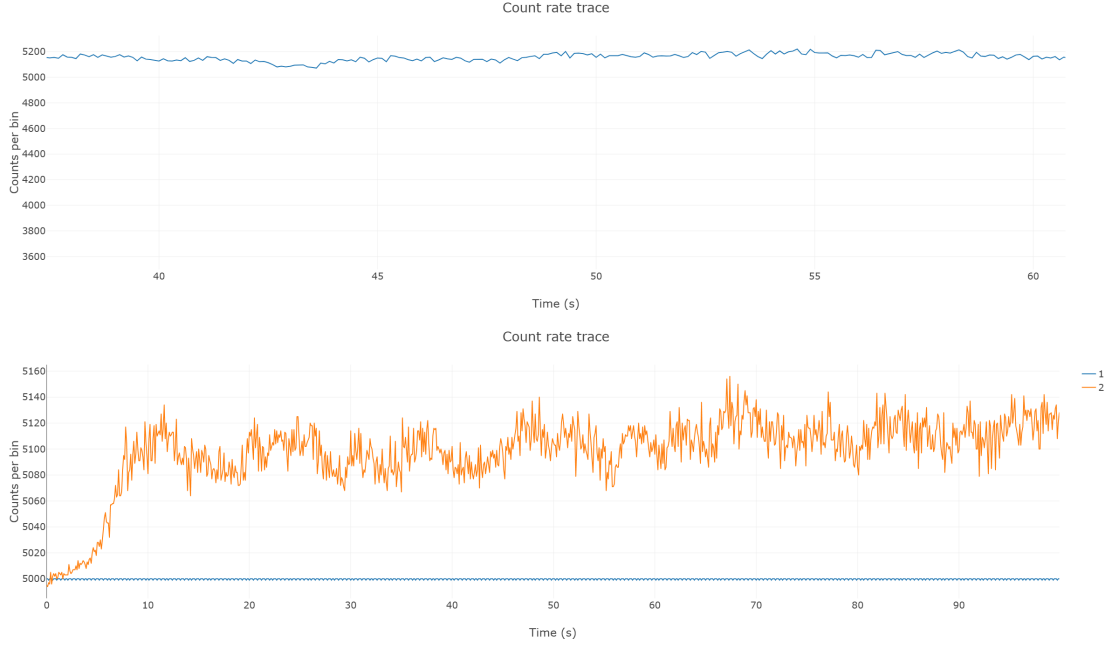


Figure 5.39: Anomalous behavior of the SPAD detector in response to a high-intensity signal. The graph shows a count that exceeds the theoretical limit of 5000 clicks, indicating a non-ideal behavior of the quenching circuit and the subsequent inability of the detector to correctly discriminate multiple pulses.

System Loss Estimation

From this experiment, we established that the condition of *at least one photon per pulse* can be satisfied with an external optical attenuation of approximately 13 dB. At this point, we compare the total losses of our system with the expected attenuation that is required to achieve the single photon regime from a classical source

Considering the quantum laser power of 19.95 mW, and a pulse duration of 2 ns, the number of photons per pulse is approximately 3.12×10^8 . To reduce this number of photons to a single photon per pulse (i.e., attenuation factor equal to 3.12×10^8), a total attenuation of approximately 85 dB is required

The intrinsic attenuation ($L_{\text{intrinsic}}$) of the Bob-Charlie path is the required value (85 dB) minus the total known losses introduced by the active components and the detector's efficiency loss. These losses are:

- Attenuation of the analog VOA (to achieve single-photon regime with 1 mW input): 13.16 dB

- Power Offset (from 1 mW to 19.95 mW): 13 dB
- Insertion Loss (IL) of the digital VOA: 3 dB
- Fixed attenuation of the digital VOA: 0 dB
- Insertion Loss (IL) of the phase modulator: 5 dB
- Detector efficiency loss: 10 dB (10 % efficiency corresponds to 10 dB loss)

NOTE Here we account for the 13 dB difference between the 19.95 mW used in this example and the 1 mW used during the VOA calibration.

Consequently, the total known loss of the system is the sum of the components, plus the 13 dB required to reach the total attenuation budget based on the 19.95 mW source:

$$L_{\text{total known}} = L_{\text{analog VOA}} + L_{\text{power offset}} + L_{\text{digital-VOA IL}} + L_{\text{digital-VOA}} + L_{\text{PM IL}} + L_{\text{detector}} \quad (5.53)$$

$$L_{\text{total known}} = 13.16 \text{ dB} + 13 \text{ dB} + 3 \text{ dB} + 0 \text{ dB} + 5 \text{ dB} + 10 \text{ dB} = 44.16 \text{ dB}$$

Therefore, the intrinsic losses of the Bob-Charlie subsystem are:

$$L_{\text{intrinsic}} = L_{\text{total required}} - L_{\text{total known}} \quad (5.54)$$

$$L_{\text{intrinsic}} = 85 \text{ dB} - 44.16 \text{ dB} = 40.84 \text{ dB}$$

This value is now consistent with the initial estimates for the Bob-Charlie path (27.1 dB + 14.5 dB = 41.6 dB), as reported in Section (5.6.4).

Chapter 6

Code and Simulation for the Asymmetric SNS-TF-QKD Protocol

In this experiment, we employed the **asymmetric version of the 4-intensity SNS TF-QKD protocol** [36], taking into account the effects of the finite size of the data exchanged between the two terminal nodes (*finite-size effects*).

The presented setup was specifically designed to adopt this protocol variant for several key reasons:

- **Security and robustness.** The SNS protocol provides unconditional security against coherent attacks, a property experimentally validated in the asymptotic regime [36]. A key feature is its significant tolerance for misalignment, relaxing the stringent precision requirements associated with single-photon interference. Simulations have demonstrated that a high key rate can be sustained even with misalignment errors approaching 35% [28].
- **Comprehensive practical effects management.** Unlike other coherent-state TF-QKD variants, the SNS protocol features a unique structure capable of effectively modeling and mitigating both the intrinsic statistical fluctuations of weak coherent states and the finite-size effects associated with the decoy states [16].
- **Optimization for asymmetric setup.** We implemented an asymmetric version of the SNS protocol because it has been demonstrated [36, 43] that with TF-QKD setups using asymmetric channels, the key rate achieved with this variant is significantly higher than that of the original SNS protocol.

This section presents the code for the realization of the chosen protocol. The objective is to demonstrate how to perform all the necessary steps for the protocol's implementation and to highlight the mathematical constraints it imposes on the selection of our setup's experimental parameters. We will also show that the symmetric case is simply a particular case of the asymmetric one, which occurs when Alice and Bob have the same channel loss towards Charlie and use identical source parameters, such as the intensities of the signal states ($s_A = s_B$) and the probability of sending coherent pulses in the Z windows ($\epsilon_A = \epsilon_B$).

In particular, to implement the required steps for this specific variant of the TF-QKD protocol and to use our setup to simulate a real key exchange, we have developed two distinct codes. The first simulates the initial part of the protocol, focusing on the preparation of the photonic pulses by Alice and Bob and their management by the respective FPGA boards. The second code, while similar in its initial phase, manages the entire analysis of the protocol, respecting all the mathematical constraints essential for its security and the subsequent post-processing phase.

6.0.1 Physical Layer Simulation and Data Generation

The first code, implemented as `FPGA_SNS_SIMULATION` and reported in additional material, is designed to prepare the raw data for the quantum communication, which constitutes the first part of the SNS-TF-QKD protocol (Section 3.5.3). In a real-world scenario, this data would be sent directly to the hardware devices in our setup, such as the FPGA boards located in the two terminal nodes (Alice and Bob). Consequently, the code's primary function is not only to simulate part of the protocol, but also to act as a command generator that accurately replicates the instructions for the physical hardware we have built.

More specifically, this code simulates the decisions made by Alice and Bob for each timeslot of the quantum communication. Based on their selection of the time window type (Signal or Decoy), the code then determines and sets the corresponding intensity and phase levels of the signal to be sent.

Process Overview and Code Structure

This process begins with defining all the protocol parameters in a series of organized dictionaries (see Table 6.1).

Then, for each timeslot, the code then generates a tuple of commands (*Intensity*, *Phase*) for each user, which represents uniquely the photonic pulse to be created. This output is crucial because it acts as the input for the intensity and phase modulators, which, at each terminal node, are controlled by their respective FPGAs. This approach enables the implementation of the SNS TF-QKD protocol's coding

procedure directly on our physical setup, precisely producing the desired quantum states.

The central part of this logic is handled by the `simulate_timeslot_preparation` function, which determines the parameters for each pulse based on a random choice of a ‘Signal’ or ‘Decoy’ window. The relevant code logic is shown in the extract that follows.

```
def simulate_timeslot_preparation(participant_name, timeslot_index):|
    """
    Simulates the preparation of a single timeslot (photonic pulse)
    by Alice or Bob, including the random choice of 'Signal' or 'Decoy'|
    window and the selection of state parameters.
    """
    # 1. Choose the window (Signal or Decoy)
    window_type = choose_window()

    # 2. Prepare the state parameters based on the chosen window
    if window_type == 'Signal':
        state_params = prepare_signal_state_params(participant_name)|
        # Public info: Bit value for raw key generation (sifting).
        public_info_for_pp = state_params['chosen_bit']
    elif window_type == 'Decoy':
        state_params = prepare_decoy_state_params(participant_name)
        # Public info: Decoy intensity for parameter estimation.
        public_info_for_pp = state_params['intensity']

    # 3. Parameters for the FPGA pulse generation
    # These values physically configure the optical modulators.
    intensity_fpga = state_params['intensity']
    total_phase_prime_fpga = state_params['phase_global_phi']
    fpga_params = (intensity_fpga, total_phase_prime_fpga)
    # ...
    return fpga_params, public_info_for_pp
```

This function produces a pair of essential values required for both the physical signal generation and the subsequent post-processing phases of the TF-QKD protocol.

In particular it returns:

1. **fpga_params (Physical Parameters for FPGA):** This tuple contains the numerical control values mapped directly to the hardware. These parameters define the physical properties of the emitted pulse and include:

- **intensity_fpga**: The quantum intensity information (mean photon number μ' or decoy states μ, ν, ω) applied to the intensity modulators.
 - **total_phase_prime_fpga**: The total quantum phase (ϕ') applied to the pulse via the phase modulators.
2. **public_info_for_pp (Public Information for Post-Processing)**: This value identifies the information to be publicly revealed later, dependent on the selected window type:
- If **window_type** is 'Signal', it contains the chosen bit value (e.g., 0 or 1) necessary for raw key generation.
 - If **window_type** is 'Decoy', it contains the selected decoy state intensity, which is fundamental for parameter estimation and security analysis.

The `main_simulation` function drives the protocol by iterating over the desired number of timeslots. Within this loop, it calls `simulate_timeslot_preparation` function for both Alice and Bob and stores the resulting hardware commands in a list. This list is then written to a text file, formatted to be usable by the hardware control software.

```
#--- MAIN FUNCTION EXECUTION ---
def main_simulation(num_timeslots):
    """
    Main loop to generate quantum states for all timeslots.
    """
    fpga_commands_list = []
    # ... Initialize other lists (e.g., simulation_results)

    for i in range(num_timeslots):
        # 1. Simulate state preparation for Alice
        alice_fpga, alice_pp_data, alice_additional_info = \
            simulate_timeslot_preparation("Alice", i)

        # 2. Simulate state preparation for Bob
        bob_fpga, bob_pp_data, bob_additional_info = \
            simulate_timeslot_preparation("Bob", i)

        # 3. Store physical parameters for the hardware (FPGA)
        fpga_commands_list.append((alice_fpga, bob_fpga))

    return simulation_results, fpga_commands_list
```

```
def save_fpga_commands_to_file(filename, fpga_commands):
    """
    Writes FPGA parameters to a file for hardware control.
    """
    for alice_params, bob_params in fpga_commands:
        alice_intensity, alice_phase = alice_params
        bob_intensity, bob_phase = bob_params

        # Format command line for the FPGA interface
```

For this experiment, the presented code is fundamental for preliminary tests and calibration. In particular, it was used to:

- Ensure that the theoretical parameters of the protocol (e.g., Intensity, phase, and sending probability) are correctly translated into commands for the hardware used.
- Generate predetermined and controlled pulse sequences to diagnose any malfunctions or inaccuracies in the chosen modulators.
- Create pulse sequences with known intensities to calibrate the detectors and verify that their responses correspond to our expected values.

Simulation Parameters and Command Generation

A simple simulation was performed to generate an example of a possible protocol output, thereby demonstrating the direct link between the protocol's parameters and the generated commands. The simulation was executed for a total of 10,000 timeslots. The protocol's configuration parameters are summarized in Table 6.1.

NOTE: In this simulation, an optimal value for M equal to 16 was assumed (refer to Section 3.4.2). The phase interval $[0, 2\pi)$ is divided into M equal slices, each with an amplitude $\Delta_k = \frac{2\pi}{M} = \frac{2\pi}{16}$. [29]

In particular, the first five timeslots of the simulation results are detailed in Table 6.2 and the correspondent FPGA commands generated are shown in Table 6.3. These commands are saved to a log file (`fpga_commands_log.txt`), which can be uploaded directly to the hardware. This functionality ensures that the theoretical parameters of the protocol are faithfully translated into practical instructions for the physical realization of the experiment.

Table 6.1: Protocol configuration parameters used in the simulation.

Category	Parameter	Value
Window Probability	Signal window probability (p_Z)	0.80
	Decoy window probability (p_X)	0.20
Alice Parameters	Signal pulse intensity (μ'_A)	0.1200
	Decoy intensity (u)	0.0100
	Decoy intensity (v)	0.0500
	Vacuum intensity (w)	0.0000
	Phase slice (Δ_k) amplitude	$\frac{2\pi}{16} \approx 0.3927$
Alice Probability	Probability to send (ϵ_A)	0.60 (Bit 1)
	Probability not to send	0.40 (Bit 0)
Bob Parameters	Signal pulse intensity (μ'_B)	0.0800
	Decoy intensity (u)	0.0031
	Decoy intensity (v)	0.0154
	Vacuum intensity (w)	0.0000
	Phase slice (Δ_k) amplitude	$\frac{2\pi}{16} \approx 0.3927$
Bob Probability	Probability to send (ϵ_B)	0.40 (Bit 0)
	Probability not to send	0.60 (Bit 1)
Asymmetry Factor	Asymmetry Factor	3.2427

Table 6.2: Protocol commands generated for the first five timeslots based on the simulation parameters.

Timeslot	Participant	Chosen Window	Sent/Not Sent	Parameters
0	Alice	Signal	Sent	(μ'_A, Δ_1)
0	Bob	Signal	Not Sent	(w, Δ_{11})
1	Alice	Decoy	-	(w, Δ_{16})
1	Bob	Decoy	-	(v, Δ_4)
2	Alice	Signal	Sent	(μ'_A, Δ_{11})
2	Bob	Decoy	-	(w, Δ_8)
3	Alice	Signal	Sent	(μ'_A, Δ_3)
3	Bob	Signal	Not Sent	(w, Δ_{16})
4	Alice	Signal	Not Sent	(w, Δ_1)
4	Bob	Signal	Not Sent	(w, Δ_{15})

Table 6.3: FPGA commands for the first five timeslots, as ready to be uploaded to hardware.

Timeslot	Alice FPGA (Intensity, Phase)	Bob FPGA (Intensity, Phase)
0	(0.12, 0.05)	(0.00, 4.06)
1	(0.00, 5.94)	(0.02, 1.23)
2	(0.12, 4.28)	(0.00, 2.96)
3	(0.12, 1.18)	(0.00, 6.02)
4	(0.00, 0.34)	(0.00, 5.88)

6.0.2 Protocol Post-Processing and Security Analysis

The second, more complex code, reported in additional material as `SNS\_SIMULATION\_COMPLETE.ipynb`, constitutes the core of the entire protocol's analysis. Unlike the initial code, which is limited to preparing the photonic pulses for transmission from the terminal nodes to the central station, this code also implements the entire post-processing phase. Specifically, it is designed to calculate the final secret key rate and evaluate the protocol's effective security under realistic conditions, accepting data generated by the previous simulation or from a real experimental run as input. The main steps of this process are described below.

Process Overview and Code Structure

The code begins by defining a series of critical parameters for the simulation and protocol analysis, organizing them into various dictionaries. These include parameters for the source, transmission channels and detectors, which are essential for a faithful reproduction of the quantum physics and system inefficiencies. They are reported in the Table 6.4:

Protocol Simulation Before the post-processing phase, the code simulates the physical behavior of the system, including data generation and detection simulation.

```
def calculate_channel_attenuation(length_km):
    """Calculates the channel attenuation factor."""
    ...
    return attenuation_factor

def simulate_charlie_click(alice_fpga_params,
                           bob_fpga_params, charlie_random_phase):
    """Simulates the arrival of photons at Charlie and their detection."""
    ...
    return click_D0_binary, click_D1_binary
```

Table 6.4: Summary of Simulation Parameters for the Asymmetric SNS-TF-QKD Protocol.

Parameter Category	Symbol	Value
Source and Channel		
Clock Rate	R_{clock}	250 MHz
Stability Overhead	τ_{stab}	1×10^{-3}
Time Slot Duration	T_{slot}	4.0 ns
Fiber Attenuation Coeff.	α	0.2 dB/km
Alice-Charlie Length	L_{AC}	20.0 km
Bob-Charlie Length	L_{BC}	20.0 km
Detector (SPAD)		
Dark Count Rate (per pulse)	P_{DC}	1×10^{-9}
Detection Efficiency	η_{det}	0.9
Detector Error Rate	e_{det}	0.005
Protocol (SNS-TF-QKD)		
Signal Intensity (Alice)	μ'_A	0.2
Signal Intensity (Bob)	μ'_B	0.2
Z-Basis Selection Prob.	p_Z	0.8
Probability to send (Alice)	ϵ_A	0.6
Probability to send (Bob)	ϵ_B	0.5
Privacy Amplification Factor	f_{Error}	1.15
Decoy State Intensities		
Decoy 1 (Alice)	μ_A	0.01
Decoy 2 (Alice)	v_A	0.05
Decoy 1 (Bob)	μ_B	≈ 0.0118
Decoy 2 (Bob)	v_B	≈ 0.0592
Decoy 3 (Vacuum)	w	0.0
Derived Asymmetry Factor		
Asymmetry Factor	A	≈ 0.845

This function returns the two binary lists, `click_D0_binary` and `click_D1_binary` that represent the time-stamped clicks registered at Charlie's two single-photon detectors (D0,D1).

State Preparation and Sifting One of the first phases of post-processing is *sifting*, which involves selecting useful data. In particular, the code prepares and filters the data to select only events where a single detection occurred.

```
def create_single_photon_detection_mask(detector0_clicks, detector1_clicks):|
    """
```

```

Creates a mask to identify timeslots with exactly one detection.
"""
# This function compares the two detector click lists to
# isolate events where exactly one click occurred.
...
return detection_mask

```

This function returns the `detection_mask`, a Boolean mask that identifies the timeslots containing exactly one detection event, which are the only valid events for key generation.

```

def create_sifting_masks_with_charlie_filter(A_choices_raw, B_choices_raw,
det_mask, A_det_mask, B_det_mask):
    """
    Classifies detection events into Z and X effective windows.
    """
    # This function uses the single-detection mask
    # and public basis choices to categorize events by basis.
    ...
    return sifting_mask_Z, sifting_mask_X

```

The function returns two Boolean masks, `sifting_mask_Z` and `sifting_mask_X`, which classify the single-detection events into the respective Z and X effective windows.

Security Parameter Estimation Using decoy state data, the code estimates security parameters, which are essential for quantifying security against an eavesdropper's attacks.

```

def H2(x):
    # Binary entropy function for privacy amplification
    ...
    return -x * math.log2(x) - (1-x) * math.log2(1-x)

```

The function returns the scalar value of the **binary entropy function** ($H_2(x)$), used to calculate the cost of privacy amplification.

```

def estimate_decoy_parameters(Ns, Nv, Nd, Na, Nc, Ns_raw, **params):
    """
    Uses the decoy state method to estimate the lower bound
    for the 1-photon yield and the upper bound for the phase error rate.
    """
    ...
    return Y1_lower_bound, e1_upper_bound

```


This function returns two crucial quantities: the lower bound for the single-photon yield $Y1_lower_bound$ and the upper bound for the phase error rate $e1_upper_bound$).

Secret Key Rate Calculation Finally the function `RdB_SNS_AOPP_TF-QKD_err` calculates the final secret key rate (R). This function accounts for all system inefficiencies, including detection errors, phase noise, and the cost of error correction. The result is a realistic estimate of the protocol's performance.

```
def RdB_SNS_AOPP_TF_QKD_err(dBloss, SigmaPhi, **var):
    # e1Upper: estimate of the phase error rate.
    ...
    # --- Final Secret Key Rate Calculation (R) ---
    R = overlap * (n1pp * (1 - H2(e1Upperpp)) - var['fError'] *
    * Nttilde * H2(Ezpp))
    return R
```

For this work, this code is indispensable for analyzing collected data and, above all, obtaining a first realistic estimate of the setup's performance. This function can be used in the future to plot key rate curves as a function of chosen experimental parameters and, consequently, to optimize the latter for the real characteristics of our setup to maximize the speed and security of key distribution using the TF-QKD protocol.

Analysis and Results from a Quasi-Realistic Model

We now present the results of a quasi-realistic simulation that provides an example of the outcomes that can be obtained using our protocol, based on the previously described code.

This simulation does not represent a definitive experimental result, as the chosen parameters, while realistic, are not the optimal ones for our specific setup. Furthermore, unlike a complete physical implementation, this simulation uses a simplified analysis for estimating crucial security parameters, such as the single-photon pair rate (n_1^{pp}) and the phase error rate (e_1^{upper}) (more details are reported in Appendix D). This simplification was a deliberate design choice to conduct a preliminary simulation while awaiting further measurements required for our actual setup (such as the exact hexadecimal levels for intensity modulation and the average number of photons per pulse). Nevertheless, in the future, the complete code can be used for an exhaustive simulation of the asymmetric SNS-TF QKD protocol, including the analysis of decoy states and finite key effects.

NOTE Despite the simplified approach, the simulation is able to incorporate

an asymmetry factor between Alice’s and Bob’s signals, making the model more faithful to the reality of the SNS-TF protocol. Moreover, it takes into account the real behavior of the sources, detectors, and channels, with their associated noise and efficiencies. Finally, it also includes a non-zero time for system re-alignment and stabilization, which is accounted for in the overall simulation time.

The temporal scale of this simulation is managed by a key parameter: *the clock rate*, which sets the duration of the QKD timeslot (*timeslot_duration_ns*). With the clock rate set at 250 MHz, the system is capable of exchanging up to 250 million signals per second. Consequently, the duration of each QKD exchange event is 4 ns (see Section 5.6.4). In particular, the entire simulation is carried out by iterating over 50 million timeslots, collectively simulating 0.2 seconds of QKD key exchange between the involved parties, in order to gather a sufficiently large statistical sample to mitigate stochastic fluctuations.

We report in Table 6.5 the key physical and protocol parameters used for the simulation. Following this, Table 6.6 summarizes the main outcomes of the analysis, highlighting the number of detected events and the estimated security parameters.

From the latter table, it can be observed that, although this simulation is based on simplified assumptions, it demonstrates that by setting optimized values for detection efficiency, attenuation, and emission probabilities—and by utilizing an appropriate model for security parameters such as n_1^{pp} and e_1^{upper} —it is possible to achieve a positive Secret Key Rate (SKR) on the order of 2×10^6 bit/s.

NOTE: This result does not represent the final optimal SKR achievable with our specific setup, but it serves as a significant example of a potential outcome under quasi-realistic conditions. Furthermore it is noteworthy that this SKR is consistent with those reported in other works on the SNS-TF-QKD protocol [45].

Table 6.5: Simulation parameters for the SNS-TF protocol.

Category	Parameter	Value	Unit
Source	<i>clockrate</i>	2.5×10^8	Hz
	<i>stab_overhead</i>	0.001	-
	<i>timeslot_duration_ns</i>	4.0	ns
Channel	<i>alpha</i>	0.2	dB/km
	<i>length_alice_charlie_km</i>	20.0	km
	<i>length_bob_charlie_km</i>	20.0	km
Detectors	<i>pDC</i>	1.0×10^{-9}	-
	<i>etadet</i>	0.9	-
	<i>detector_error_rate</i>	0.005	-
Protocol (SNS)	<i>pZ_SNS</i>	0.8	-
	<i>pX_SNS</i>	0.2	-
	ϵ_A	0.6	-
	ϵ_B	0.5	-
	<i>intensity_A_signal_pulse</i>	0.2	photons/pulse
	<i>intensity_B_signal_pulse</i>	0.2	photons/pulse
Decoy and Asymmetry	<i>Asymmetric_Factor</i>	1.5	-
	<i>intensity_A_u_decoy</i>	0.01	photons/pulse
	<i>intensity_A_v_decoy</i>	0.05	photons/pulse
	<i>intensity_A_w_decoy</i>	0.0	photons/pulse
	<i>intensity_B_u_decoy</i>	0.0067	photons/pulse
	<i>intensity_B_v_decoy</i>	0.0333	photons/pulse
	<i>intensity_B_w_decoy</i>	0.0	photons/pulse

Table 6.6: Results of the SNS-TF protocol simulation.

Description	Value	Unit	Symbol
Total timeslots simulated	50,000,000	-	N_{total}
Events with single click	5,418,436	-	N_{single}
Events in Z window (signal)	4,273,204	-	N_Z
Events in X window (decoy)	11,896	-	N_X
Number of paired events	4,285,100	-	$N_{\tilde{t}}$
Paired events error rate	0.264	-	E_{zpp}
Lower bound for single-photon yield	0.9972	-	Y_1^{lower}
Upper bound for phase error rate	0.0075	-	e_1^{upper}
Secret key rate (SKR)	2,228,434.93	bit/s	R

Chapter 7

Conclusions and Future Developments

The primary objective of this thesis was the design, development, and characterization of a portable, dual-band stabilized setup for the TF-QKD protocol. The system was optimized for integration into existing long-haul telecommunication infrastructures. This work successfully demonstrated the robust performance and technical feasibility of the system, emphasizing portability through the strategic use of commercial components, such as single-photon avalanche diodes and narrow-bandwidth lasers. The main contributions include comprehensive hardware development, detailed system characterization, and rigorous theoretical security analysis.

This experimental setup exploited the most effective topology yet reported for stabilizing the relative phase between the twin fields. Specifically, we adopted the independent laser configuration with dual-band stabilization, demonstrated as the most promising approach for creating a compact and transportable system suitable for the real-world application of TF-QKD in modern telecommunication networks. This robust stabilization scheme, combined with wavelength multiplexing, successfully overcame the intrinsic critical limitations in previous TF-QKD implementations.

Crucial elements of our setup were the developed phase and temporal alignment routines, which demonstrated essential stability for the protocol. In particular, the fast phase drift compensation maintained a QBER below 0.5% for over 4 seconds in a *proof-of-concept* experiment, even under an estimated total terminal losses of 65 dB. This result overcomes the main limitation that has previously limited TF-QKD application in a real-world setting. Furthermore, slow phase drift was managed effectively, showing that re-alignment windows of just 100 ms are sufficient

for phase compensation and a precise temporal re-alignment system was developed—supported by an automated cross-correlation routine—capable of exploiting coarse and fine delay mechanisms to compensate for fiber length differences up to 50 km and minimal variations of 250 ps. The system’s tolerance to phase drift, thanks to the balance between open-channel stabilization and active control feedbacks, proves its suitability for future integration into existing telecommunication infrastructures, even for optical distances greater than 200 km.

Detailed device characterization established the optimal operating conditions and system vulnerabilities. The Raman noise analysis demonstrated the feasibility of TF-QKD communication over a total distance of 200 km, as the Raman noise (≈ 50 counts in 100 ms) was found to be tolerable when compared with the intrinsic detector background noise (≈ 20 counts in 100 ms). Furthermore, it was demonstrated that, for a realistic, long-distance implementation of TF-QKD, the system must guarantee a total ER of at least 40 dB, which was satisfied by employing two intensity modulators connected in cascade, each characterized to ensure a minimum ER of 20 dB. Comprehensive node loss characterization determined the estimated losses of 26.4 dB at Alice’s node, 27.1 dB at Bob’s node, and an additional 14.5 dB within the central Charlie node. We also appropriately analyzed the performance of other essential devices of the setup, such as the amplifiers, VOAs and lasers, to better understand their real characteristics and the consequent vulnerabilities of the proposed scheme.

Theoretically, we evaluated the asymmetric Sending-Not-Sending (SNS) variant, confirming its superior robustness and SKR compared to the symmetric approach. This variant was identified and described as the most robust and secure for our setup. A complete simulation model was implemented, which validated the results by obtaining an SKR compatible with reference implementations in the 60 dB attenuation regime. The two codes we created were fundamental to this process: the first allowed us to generate the commands for the preparation of the quantum pulses, while the second enable for a detailed analysis of the protocol and the estimation of the secret key rate. Specifically, FPGA control codes were developed to manage the intensity modulation of both signal and decoy states and the optimal phase randomization ($M_{opt} = 16$ levels).

Despite this strong foundation, the realization of a fully operational, long-distance TF-QKD system requires further development to maximize performance and reliability. In the future, the final optimization of protocol parameters (such as the number of photons per pulse and the sending and non-sending probabilities) must be calibrated using data collected from real experiments, for which the implemented codes will serve as essential resources. It will also be crucial to appropriately evaluate the characteristics of real fibers, such as lengths, losses, and environmental noise,

which can significantly increase the QBER and to optimize detector characteristics to minimize background noise. Furthermore, the phase compensation system must be validated over longer distances and the temporal re-alignment procedure should be integrated with advanced timing protocols, such as White Rabbit, to achieve nanosecond-level synchronization. Finally, the feedback for the polarization control must be fully automated to mitigate polarization mismatch. The phase drift compensation technique could also be adapted for quantum experiments in free space.

In conclusion, this work constitutes a significant step toward the large-scale integration of the TF-QKD protocol, establishing the necessary foundation to overcome the distance limitations of current QKD systems and paving the way for a global quantum communication network without relying on either trusted nodes or quantum repeaters. In particular this work lays the foundation for a future implementation of the protocol that can be further enhanced and integrated into phase-based quantum communication networks.

Appendix A

Practical Procedure for Phase-Locking

Quantum Laser Phase-Locking

In this section, we provide a detailed description of the laboratory steps required to phase-lock the quantum lasers of Alice and Bob to their respective sensing lasers.

- 1. Preparation and Open-Loop Initialization** The initial step is to bring the quantum laser's (RIO) frequency to within approximately 10 MHz of the $m = 4$ sideband of the sensing laser (Fig.[(5.5)]). This is accomplished by manually adjusting the quantum laser's temperature. The software interface allows for precise control of both temperature and current, which directly influence the emission frequency of the semiconductor laser (Fig.[(B.2)]). This step is critical because the beat note frequency must fall within the band-pass filter bandwidth to be detected, a prerequisite for the PLL to function correctly.
- 2. Loop Closure and PID Activation** Once the beat note is within few MHz from the nominal locking frequency of 10 MHz, the PLL loop is closed, and the Proportional-Integral-Derivative (PID) controller is activated. Initially, only the proportional (P) component is enabled to provide an immediate correction. A manual offset can be applied to further optimize the signal. After a stable lock at 10 MHz is achieved, the integral (I) component of the PID is enabled and ensuring robust lock also on long timescales eliminating steady-state error. Generally, the PID controller acts on the quantum laser's drive current.
- 3. Slow-Loop Implementation** The slow loop is an auxiliary control loop that is activated as the final phase of the locking procedure. Its function is to prevent the PID controller from having to apply excessive current corrections

to compensate for slow frequency drifts, such as those caused by thermal variations. By activating the slow loop, the laser's temperature is regulated to keep the drive current within an optimal operating range. This ensures that the fast loop (which controls the current) can operate at maximum effectiveness, handling only rapid corrections.

Sensing Laser Phase-Locking

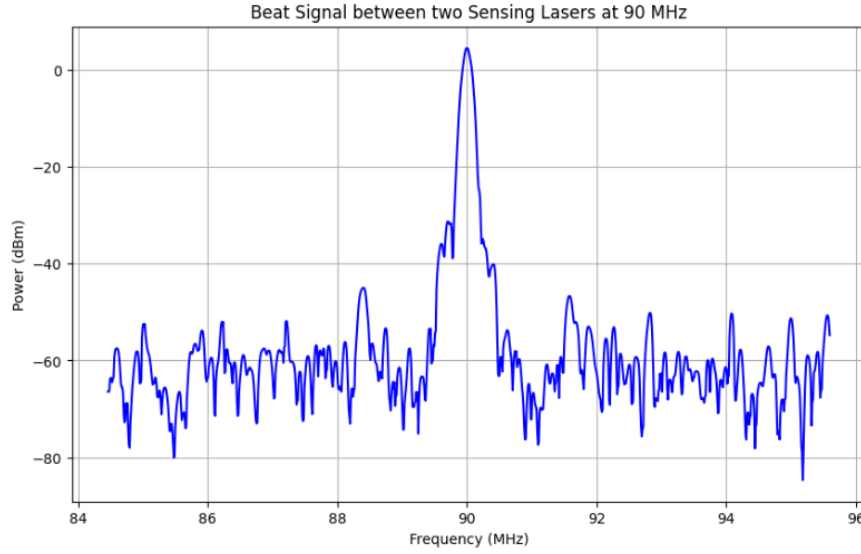


Figure A.1: Spectrum of the beat signal between the *sensing* lasers of Alice and Bob at 90 MHz.

In this section, the practical laboratory steps required to phase-lock Alice's and Bob's sensing lasers are outlined. This procedure, corresponding to the third PLL of our system (detailed in Fig.[5.8]), is crucial for correcting phase differences that may arise from the intrinsic noise of the two quantum sources and subsequent phase drift accumulated along the optical fibers connecting the terminal nodes to Charlie.

The process is divided into two main phases: *Open-Loop Initialization* and *Closed-Loop Operation*.

1. **Open-Loop Preparation and Beat Note Visualization** The procedure begins with the control loop in an open configuration, with the primary goal of visualizing the beat note frequency between Alice's and Bob's sensing lasers (Fig.[A.1]). This is achieved by connecting the output of the photodiode (PD), which combines the two laser signals, to a spectrum analyzer.

The spectrum analyzer displays the beat note signal, which is equal to the sum of the AOM1-induced phase shift (Δ) and the frequency shift caused by fiber noise. Subsequently, a manual offset voltage is applied to the VCO that drives AOM1 to bring the beat note within few MHz from the locking frequency. This offset is essential for precisely setting the beat note frequency. As described in the theoretical scheme (Section 5.3.2), in our setup, the offset is adjusted to center the beat note signal at approximately 90 MHz on the spectrum analyzer, a frequency that corresponds to the total shift induced by AOM1.

This step ensures that the beat note falls within the operating range of the control electronics, preparing the system for the next phase.

2. **Closed-Loop Operation and Active Phase Locking** Once the beat note is stably positioned at 90 MHz (Fig.[A.1]), the control loop is closed. The VCO, in turn, converts this voltage into a correction frequency that is sent to AOM1. Since AOM1 acts as the actuator in this feedback loop, it applies this correction frequency—with the opposite sign—to Alice’s multiplexed signal. This dual-pass frequency shift allows the system to actively cancel out the relative phase difference and the accumulated fiber noise between the two sensing lasers.

The continuous feedback loop ensures that the absolute frequency difference between the two sensing lasers remains stable and equal to the offset frequency provided by the VCO, thus successfully phase-locking them. This active stabilization in the Charlie node enables the system to maintain the coherence of the quantum lasers over long distances, correcting for environmental and the lasers frequency fluctuations.

Appendix B

Experimental parameters and technical details

Figures [B.2, B.1] show the control panels of the sensing and quantum laser with the setting point parameters that correspond to the designed frequency offset, and optimal loop gains settings.

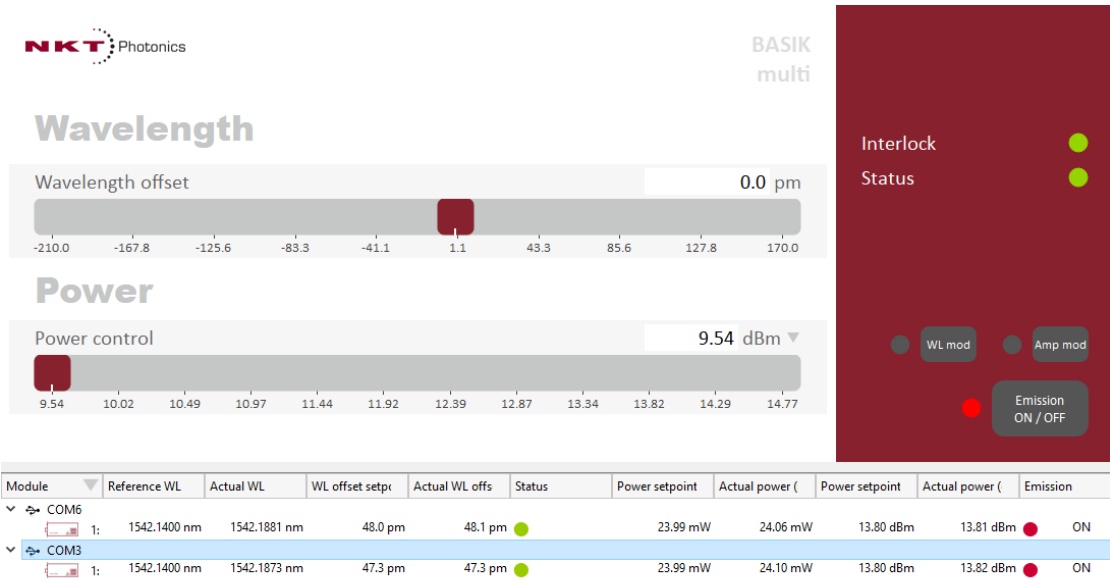


Figure B.1: Software interface for controlling the sensing lasers (NKT) of Alice and Bob

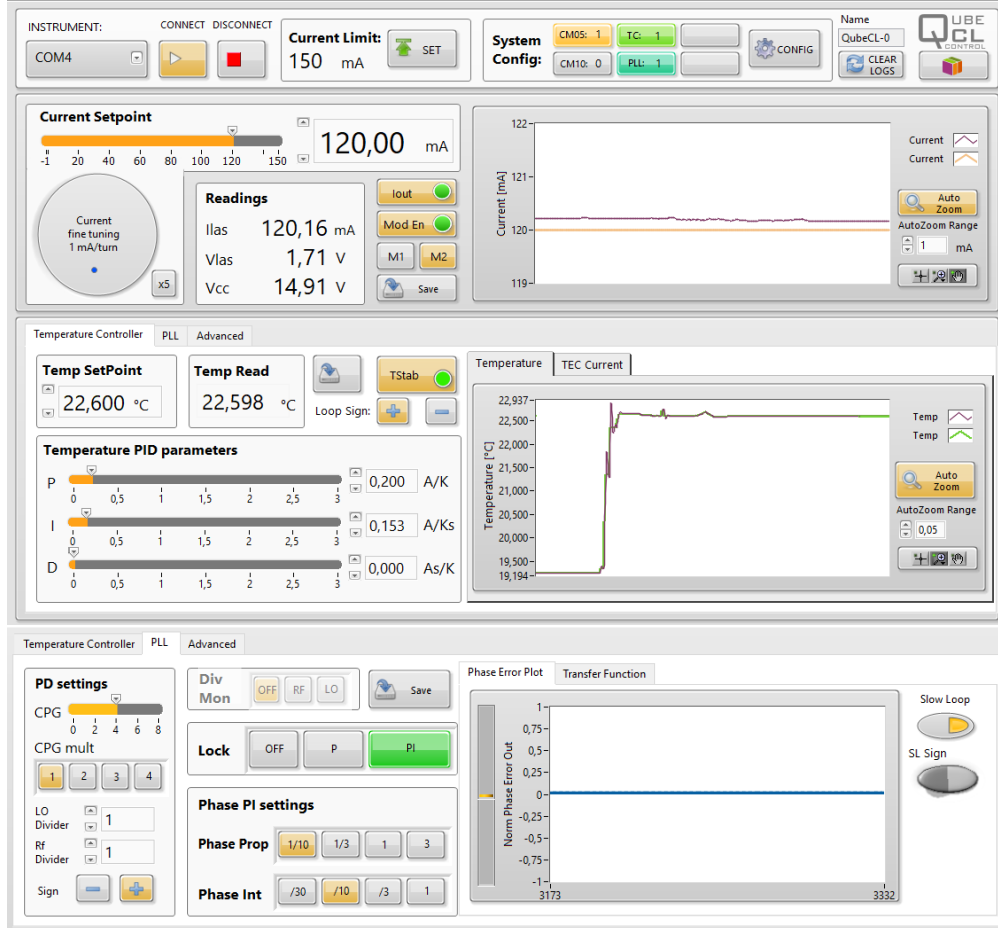


Figure B.2: Software interface for controlling the quantum laser (RIO). The top panel shows the main parameters for current and temperature control, while the bottom panel shows the PID controller settings and the real-time feedback loop status.

Appendix C

Technical Steps for Delay Setting

To delve into the technical details of the temporal alignment procedure, let's examine the code implemented to compensate for the delay using the coarse and fine delay settings via the FPGA software board (see *alice_bob_delays* code in Additional material). The delay corrections are implemented through a two-phase strategy:

1. **An initial calibration** that compensates for large-scale path differences.
2. **A continuous feedback operation** to maintain alignment.

Coarse Delay Setting The coarse delay is the primary step in temporal alignment, designed to compensate for the time-of-flight differences between pulses from Alice and Bob that are primarily caused by unequal fiber lengths. This procedure can be executed preliminarily as a "pre-communication calibration protocol" to maximize initial interference visibility, or it can be performed periodically whenever the physical setup changes significantly.

Within our setup, the *alice_bob_delays* FPGA code is responsible for setting the control registers for the **coarse delay**. Since each coarse step is equal to 4 ns, this enables large-scale compensation. This entire process is supported by an external time base from a function generator, which simulates the White Rabbit protocol and provides the essential nanosecond synchronization for our experiment.

The configuration of the channels in reference mode is illustrated by the following code lines:

```
base_pps_gen.write(4, 165000)
base_pps_gen.write(0, 1)
```

These commands activate the generation of a hardware-level synchronization signal, which serves as a reference point for subsequent adjustments. Subsequently, the predetermined delay values for Alice and Bob (`delay_alice` and `delay_bob`) are written to the FPGA registers, applying the coarse delay to all modulators of each node.

The following snippet from the code demonstrates this step:

```
# Coarse delays for Alice's modulators
base_gpio_ctr.write(df_a_IM1 * 4, 0 + delay_alice)
base_gpio_ctr.write(df_a_IM2 * 4, 0 + delay_alice)
base_gpio_ctr.write(df_a_PM * 4, 0 + delay_alice)
# Coarse delays for Bob's modulators
base_gpio_ctr.write(df_b_IM1 * 4, 0 + delay_bob)
base_gpio_ctr.write(df_b_IM2 * 4, 0 + delay_bob)
base_gpio_ctr.write(df_b_PM * 4, 0 + delay_bob)
```

This step ensures that the pulses from Alice and Bob arrive at Charlie with a sufficiently precise temporal alignment to allow for subsequent adjustments.

Fine Delay Setting After applying the coarse delay, the process moves to sub-nanometric precision alignment to align the modulation windows between Alice and Bob.

This phase of the realignment is crucial not only to ensure that the signals arrive at Charlie with a perfect temporal overlap within a 2 ns modulation window, but also to ensure that the modulators within each node are perfectly synchronized with each other. It is assumed that this latter alignment is performed only once at the beginning using a classical setup and an oscilloscope to determine the delay, while conversely, to finely align the modulation windows between Alice and Bob, this phase is performed periodically using the same *alice_bob_delays* FPGA code, always after the coarse delay has been set. To be precise, as each fine step is equal to 250 ps, this allows for highly precise adjustments.

The following lines of code write the previously determined fine delay values to their respective control registers:

```
# Fine delays for Alice's modulators
base_gpio_ctr.write(df_a_IM1 * 4, 0 + delay_alice)
base_gpio_ctr.write(df_a_IM2 * 4, 12 + delay_alice)
base_gpio_ctr.write(df_a_PM * 4, 5 + delay_alice)
# Fine delays for Bob's modulators
base_gpio_ctr.write(df_b_IM1 * 4, 24 + delay_bob)
base_gpio_ctr.write(df_b_IM2 * 4, 2 + delay_bob)
```

```
base_gpio_ctr.write(df_b_PM * 4, 6 + delay_bob)
```

These values, such as `df_a_IM1` and `df_a_IM2`, serve as specific offsets for each modulator, enabling extremely precise alignment. This configuration is essential for optimizing pulse overlap, minimizing jitter, and ensuring the stability of the system for the TF-QKD protocol.

Combined Result In conclusion, the combination of coarse and fine delay creates a robust and efficient calibration system. The coarse delay handles large-scale compensations, while the fine delay makes the necessary refinements for quantum precision.

Appendix D

Simulation Simplifying Hypotheses

This section outlines the key assumptions and approximations adopted in the preliminary simulation of the SNS-TF QKD protocol, as detailed in Section 6.7. These choices were made with the dual objective of confirming the algorithm’s viability in a near-ideal setting and securing a positive result for the Secret Key Rate before a comprehensive experimental characterization of our setup could be performed.

The formulas used to compute the lower bound of the single-photon yield and the upper bound of the phase error rate were adapted to serve the purpose of this simulation.

Lower Bound for Single-Photon Yield (Y_1^{lower})

The formula $Y_1^{\text{lower}} = \frac{N_Z - N_X}{N_Z}$ if $N_Z > 0$ otherwise 0 represents a simplified model for the decoy state method. The purpose of this method is to approximate the single-photon yield by filtering out events originating from non-ideal states (e.g., multi-photons). Within this simplified framework, the count from the ‘Decoy’ basis (N_X) is treated as a measure of the noise present in the ‘Signal’ basis (N_Z) [38]. By subtracting N_X from N_Z , we derive a preliminary estimate of the single-photon contribution, which is sufficient for the SKR calculation.

Upper Bound for Phase Error Rate (e_1^{upper})

The assignment $e_1^{\text{upper}} = \text{detector_error_rate} \times 1.5$ serves as a practical assumption for a proof-of-concept simulation. In a real-world system, phase errors result from a multitude of factors, including modulator imperfections, thermal instabilities, and

various other noise sources. For the porpoise of this simulation, these contributions are collectively represented by an aggregated parameter: the *detector_error_rate*. This approach avoid the need for a detailed statistical analysis of phase errors from the X basis, while still demonstrating the protocol's robustness in a low-noise environment.

Adjustment to the Paired Events Error Rate

Base on these simplifications, a key adjustment was made to the calculation of the paired events error rate (E_{zpp}) to guarantee a consistent outcome for the SKR. The specific value was chosen to be small and representative of a realistic scenario.

- **Formula:** $E_{zpp} = \text{var}[\text{'detector_error_rate'}] * 1.5$

This modification ensures that E_{zpp} is derived directly from the `detector_error_rate`. This prevents the data correction costs from exceeding acceptable thresholds.

Appendix E

Comparison of typical properties of APDs and SNSPD

Table E.1: A summary table outlining the key technological differences between APD (Avalanche Photodiode) and SNSPD (Superconducting Nanowire Single-Photon Detector) detectors [45].

Property	APDs	SNSPDs
Detection efficiency (%)	10–20	60–80
Dark counts	Hundreds of hertz	A few hertz
Afterpulse (%)	About 3	0
Size (U)	<1	>10
Cost (\$)	1,000s	100,000s
Cooling system	Thermoelectric (20 to -60)°C	Cryogenic (-270 to -272)°C

Within the field of quantum telecommunications, the primary solutions for single-photon detection are Superconducting Nanowire Single-Photon Detectors (SNSPDs) and Single-Photon Avalanche Photodiodes (SPADs), typically fabricated from InGaAs/InP. It is crucial to evaluate the characteristics and performance of these detectors, as they directly influence the type of system that can be deployed and the maximum communication distance achievable by a given protocol.

For a detailed technical comparison, Table (E.1) highlights the key distinctions between these two detector technologies [45]. The analysis yields a few critical

insights.

SNSPDs are currently considered the most sophisticated and high-performing detectors available. They are characterized by high detection efficiency (60-80%) and an exceptionally low dark count rate (just a few Hz). However, their major drawbacks are the high cost and the absolute necessity of a cryogenic cooling system, which operates at an extreme temperature of approximately -270°C . These requirements make SNSPDs bulky and limit their practicality and scalability for field applications.

Conversely, SPADs offer a much more practical and accessible solution. While their performance is more modest, with a detection efficiency of 10-20% and a higher dark count rate (hundreds of Hz), they operate using more manageable thermoelectric cooling (-20°C to -60°C). This less demanding cooling requirement, combined with a cost that is one or two orders of magnitude lower and a significantly smaller form factor (typically $< 1U$ compared to the $> 10U$ required for SNSPDs), makes them a more suitable choice for commercial and on-site deployments. In essence, SPADs strike an effective balance between performance, cost, and scalability [12, 45].

Bibliography

- [1] Christof Paar and Jan Pelzl. *Understanding Cryptography: A Textbook for Students and Practitioners*. Springer, 2010 (cit. on pp. 1–11, 13, 14).
- [2] Susan Loepp and William K. Wootters. *Protecting Information: From Classical Error Correction to Quantum Cryptography*. Cambridge University Press, 2006 (cit. on pp. 9, 20, 27, 70).
- [3] Claude E. Shannon. «Communication Theory of Secrecy Systems». In: *The Bell System Technical Journal* 28.4 (1949), pp. 656–715. DOI: 10.1002/j.1538-7305.1949.tb00924.x (cit. on p. 12).
- [4] Jonathan P. Dowling and Gerard J. Milburn. «Quantum Technology: The Second Quantum Revolution». In: *The Physics of Quantum Information*. Ed. by Dirk Bouwmeester, Artur Ekert, and Anton Zeilinger. Springer, 2003, pp. 127–142 (cit. on p. 16).
- [5] Robert S. Sutor. *Dancing with Qubits*. 2nd ed. Packt Publishing, 2020 (cit. on pp. 17, 22, 29, 32).
- [6] Fabrice Boudot, Pierrick Gaudry, Aurore Guillevic, Nadia Heninger, Emmanuel Thomé, and Paul Zimmermann. «Factorization of RSA-250». In: *Cryptology ePrint Archive, Report 2020/058* (2020) (cit. on p. 18).
- [7] National Security Agency. *Commercial National Security Algorithm Suite*. Fact Sheet. Aug. 2015. URL: <https://www.nsa.gov/Press-Room/Press-Releases-Statements/Press-Release-View/Article/1621942/commercial-national-security-algorithm-suite/> (cit. on p. 18).
- [8] S. Pirandola et al. «Advances in Quantum Cryptography». In: *Advances in Optics and Photonics* 12.4 (2020), pp. 1–61. DOI: 10.1364/AOP.381387 (cit. on pp. 19, 21, 27–29, 32, 35, 37–39, 42–44, 49, 52, 54–56, 59, 66, 68, 70, 84).
- [9] Charles H. Bennett and Gilles Brassard. «Quantum cryptography: Public key distribution and coin tossing». In: *Proceedings of the IEEE International Conference on Computers, Systems and Signal Processing*. 1984, pp. 175–179 (cit. on pp. 20, 21, 32).

- [10] Bruce Schneier. *Applied Cryptography: Protocols, Algorithms, and Source Code in C*. 2nd ed. New York, NY: John Wiley & Sons, 1996 (cit. on p. 21).
- [11] S. P. Kish, P. J. Gleeson, A. Walsh, P. K. Lam, and S. M. Assad. «Comparison of Discrete Variable and Continuous Variable Quantum Key Distribution Protocols with Phase Noise in the Thermal-Loss Channel». In: *Physical Review A* 93.3 (2016), p. 032320. DOI: 10.1103/PhysRevA.93.032320 (cit. on pp. 21, 43).
- [12] Davide Calonico. *Laboratories of Quantum Technologies*. Course at LM Quantum Engineering, A.A. 2024/2025. 2024 (cit. on pp. 21, 28, 30, 38, 43, 73, 80, 81, 148, 164, 193).
- [13] Michael A. Nielsen and Isaac L. Chuang. *Quantum Computation and Quantum Information*. Cambridge University Press, 2010 (cit. on pp. 22, 25, 29).
- [14] Edo Waks, Assaf Zeevi, and Yoshihisa Yamamoto. «Security of quantum key distribution with entangled photons against individual attacks». In: *Physical Review A* 65.5 (2002), p. 052310. DOI: 10.1103/PhysRevA.65.052310 (cit. on pp. 22, 32, 39).
- [15] R. Proietti. *Introduction to Quantum Communication Systems and Networks*. Course at Politecnico di Torino, A.A. 2024/2025. 2024 (cit. on pp. 24, 26, 33, 74).
- [16] G. Bertaina and M. Boffi. «Recent advances in long-distance TF-QKD». In: *Journal of Modern Optics* 60.1 (2024), pp. 100–115 (cit. on pp. 28, 29, 41, 50, 53, 57–60, 74, 76–81, 83, 84, 88, 89, 91, 167).
- [17] Charles H. Bennett. «Quantum cryptography using any two nonorthogonal states». In: *Physical Review Letters* 68 (1992), p. 3121 (cit. on p. 32).
- [18] Ivan B. An, B. C. Sanders, and L. L. Zhang. «Unambiguous State Discrimination with Mixed States». In: *Physical Review A* 66.3 (2002), p. 032306 (cit. on p. 33).
- [19] Artur K. Ekert. «Quantum cryptography based on Bell’s theorem». In: *Physical Review Letters* 67 (1991), p. 661 (cit. on p. 33).
- [20] John F. Clauser, Michael A. Horne, Abner Shimony, and Richard A. Holt. «Proposed Experiment to Test Local Hidden-Variable Theories». In: *Physical Review Letters* 23 (1969), pp. 880–884 (cit. on p. 33).
- [21] John S. Bell. «On the Einstein Podolsky Rosen paradox». In: *Physics* 1 (1964), p. 195 (cit. on pp. 33, 35, 36, 51).
- [22] Charles H. Bennett, Gilles Brassard, and N. David Mermin. «Quantum cryptography without Bell’s theorem». In: *Physical Review Letters* 68 (1992), p. 557 (cit. on p. 34).

- [23] Llu'is Masanes, St'e'phane Pironio, and Antonio Ac'ın. «Secure device-independent quantum key distribution». In: *Nature Physics* 5 (2009), pp. 41–45 (cit. on p. 35).
- [24] Morteza Moradi, Maryam Afsary, Piotr Mironowicz, Enky Oudot, and Magdalena Stobi'nska. «Long-range photonic device-independent quantum key distribution using SPDC sources and linear optics». In: *arXiv preprint arXiv:2407.15174* (2024) (cit. on p. 35).
- [25] Monika E Mycroft, Thomas McDermott, Adam Buraczewski, and Magdalena Stobi'nska. «Proposal for the distribution of multiphoton entanglement with optimal rate-distance scaling». In: *arXiv preprint arXiv:2407.13606* (2024) (cit. on pp. 35, 36).
- [26] John G Eberhard. «Bell's theorem and the EPR paradox». In: *Foundations of Physics* 23.6 (1993), pp. 881–893 (cit. on p. 35).
- [27] Hoi-Kwong Lo, Marcos Curty, and Bing Qi. «Measurement-Device-Independent Quantum Key Distribution». In: *Physical Review Letters* 108 (2012), p. 130503 (cit. on pp. 36, 43, 44, 49).
- [28] Xiang-Bin Wang. «A review on the decoy-state method for practical quantum key distribution». In: *arXiv preprint arXiv:quant-ph/0509084* (2005). v3 (cit. on pp. 36, 39, 43, 44, 53, 60, 61, 71, 167).
- [29] Marco Lucamarini, Z. L. Yuan, J. F. Dynes, and A. J. Shields. «Overcoming the rate–distance limit of quantum key distribution without quantum repeaters». In: *Nature* 557.7705 (2018), pp. 400–403. DOI: 10.1038/s41586-018-0211-1 (cit. on pp. 37, 43, 44, 46–48, 50, 52, 53, 55, 73, 171).
- [30] Zhen-Qiang Yin, Sheng-Kai Liao, Yan-Hui Zhao, Hai-Han Wu, Jian-Jun Liu, Chao-Yong Lu, and Jian-Wei Pan. «Twin-field quantum key distribution with unbalanced sources». In: *Physical Review A* 100.3 (2019), p. 032320. DOI: 10.1103/PhysRevA.100.032320 (cit. on p. 42).
- [31] Zhen-Qiang Yin and et al. «Real-world twin-field quantum key distribution over 509 km fiber». In: *Physical Review Letters* 124.1 (2020), p. 010506. DOI: 10.1103/PhysRevLett.124.010506 (cit. on pp. 42, 43, 45, 51, 52, 56).
- [32] Gianluca Guidi et al. «Continuous-variable quantum key distribution: from a proof-of-principle experiment to a field-test implementation». In: *2018 European Conference on Optical Communication (ECOC)*. 2018, pp. 1–3. DOI: 10.1109/ECOC.2018.8530064 (cit. on p. 43).
- [33] Marcos Curty, Koji Azuma, and Hoi-Kwong Lo. «Simple security proof of twin-field type quantum key distribution protocol». In: *npj Quantum Information* 5.1 (2019), pp. 1–7. DOI: 10.1038/s41534-019-0158-9 (cit. on p. 45).

- [34] Xiongfeng Ma, Hong-Wei Jiang, Zong-Wen Sun, and Zheng-Wei Zhou. «Practical security analysis of sending-or-not-sending twin-field quantum key distribution». In: *Quantum Science and Technology* 4.4 (2019), p. 045012 (cit. on pp. 56, 58, 60, 61).
- [35] Charles C W Lim, Marcos Curty, and Christoph Leinfellner. «Coherent-state attacks on a twin-field quantum key distribution system». In: *Physical Review A* 98.2 (2018), p. 022312 (cit. on pp. 56, 58).
- [36] Zhengyuan Xiao, Jiageng Chen, Jiazhen Ji, Qingwen Liu, and Zuyuan H. «Complete Phase Noise Compensation for 50 km DAS with 100 kHz linewidth ITLA». In: *Journal of Lightwave Technology* 40.17 (2022), pp. 5778–5784 (cit. on pp. 60, 62–66, 69, 71, 72, 167).
- [37] Xiao-Long Hu, Cong Jiang, Zong-Wen Yu, and Xiang-Bin Wang. «Sending-or-not-sending twin-field protocol for quantum key distribution with asymmetric source parameters». In: *Physical Review A* 99.3 (2019), p. 032331. DOI: 10.1103/PhysRevA.99.032331 (cit. on p. 60).
- [38] Cong Jiang, Zong-Wen Yu, Xiao-Long Hu, and Xiang-Bin Wang. «Unconditional security of sending or not sending twin-field quantum key distribution with finite pulses». In: *Physical Review A* 100.4 (2019), p. 042318. DOI: 10.1103/PhysRevA.100.042318 (cit. on pp. 61, 65–68, 190).
- [39] Lai Zhou, Jinping Lin, Yumang Jing, and Zhiliang Yuan. «Twin-field quantum key distribution without optical frequency dissemination». In: *npj Quantum Information* 7.1 (2021), pp. 1–7. DOI: 10.1038/s41534-021-00366-z (cit. on pp. 64, 75, 92–94, 99).
- [40] Hoi-Kwong Lo, Marcos Curty, and Bing Qi. «Practical quantum key distribution». In: *Science Advances* 4.6 (2018) (cit. on p. 69).
- [41] Cecilia Clivati et al. «Coherent phase transfer for real-world twin-field quantum key distribution». In: *Nature Communications* 12.1 (2021), pp. 1–9. DOI: 10.1038/s41467-021-27808-1 (cit. on pp. 74, 81–85, 87, 95, 98, 117, 118, 122).
- [42] Alan D Kersey, Michael A Davis, Helen J Patrick, Michel LeBlanc, K P Koo, A M Glass, M A Putnam, and J S Sirkis. «Fiber-optic sensor technology and applications». In: *Journal of Lightwave Technology* 15.8 (1997), pp. 1442–1463 (cit. on p. 75).
- [43] S. Wang et al. «Phase-referencing technique for twin-field quantum key distribution». In: *Physical Review X* 9.2 (2019). URL: <https://journals.aps.org/prx/abstract/10.1103/PhysRevX.9.021046> (cit. on pp. 78, 80, 81, 167).

- [44] Gianluca Bertaina, Cecilia Clivati, Davide Calonico, Marco Genovese, and Alberto Vallone. «Phase Noise in Real-World Twin-Field Quantum Key Distribution». In: *arXiv preprint arXiv:2203.01321* (2022). DOI: 10.48550/arXiv.2203.01321 (cit. on p. 82).
- [45] Mirko Pittaluga et al. «Long-distance coherent quantum communications in deployed telecom networks». In: *Nature Communications* 12.1 (2021), pp. 1–8. DOI: 10.1038/s41467-021-26012-y (cit. on pp. 87, 177, 192, 193).
- [46] European Commission. *European Quantum Communication Infrastructure (EuroQCI)*. 2024. URL: <https://digital-strategy.ec.europa.eu/en/policies/european-quantum-communication-infrastructure-euroqci> (cit. on p. 88).
- [47] Bahaa E. A. Saleh and Malvin Carl Teich. *Fundamentals of Photonics*. John Wiley Sons, 1991 (cit. on pp. 102, 103).
- [48] William K. Burns and Theodore G. Giallorenzi. «Noise techniques for the analysis of the intrinsic 0 to 20kHz line broadening of CW GaAs lasers». In: *Journal of Applied Physics* 46.5 (1975), pp. 1895–1897 (cit. on p. 120).
- [49] O. Rey, G. Daniluk, M. G. Barón, P. H. S. Diniz, L. S. Cunha, L. A. S. B. Junior, K. W. J. B. de Oliveira, and L. R. M. F. de Oliveira. «White Rabbit: A sub-nanosecond timing solution for large-scale physics experiments». In: *Nuclear Instruments and Methods in Physics Research Section A: Accelerators, Spectrometers, Detectors and Associated Equipment* 640.1 (2011), pp. 29–38. ISSN: 0168-9002. DOI: 10.1016/j.nima.2010.11.096 (cit. on pp. 129, 131).
- [50] Swabian Instruments. *Time Tagger User Manual*. Swabian Instruments. 2025 (cit. on p. 133).
- [51] Robert H. Hadfield. «Single-photon detectors for quantum information applications». In: *Nature Photonics* 3.12 (2009), pp. 696–705 (cit. on pp. 145, 146).
- [52] JENOPTIK Optical Systems GmbH. *Integrated-optical modulators: Technical information and instructions for use*. JENOPTIK. 2025 (cit. on pp. 154, 155).

Acknowledgements

Questa tesi è dedicata a tutte le persone che non hanno mai smesso di credere in me e che mi sono state accanto fino a questo momento.

Ovviamente, com'è ormai ben noto, senza le schiscette di mamma Bea e il motorino di papà Guido non sarei mai arrivato fino a questo punto.

Grazie per il vostro amore incondizionato e per il costante supporto in ogni passo di questo percorso.

Infine, come la leggenda insegna:

"Il mio segreto... è che non c'è un ingrediente segreto! Sei solo tu!"

(Cit. La Rivelazione di Mr. Ping sugli spaghetti, Kung Fu Panda)